

POLSKA AKADEMIA NAUK
KOMITET ELEKTRONIKI I TELEKOMUNIKACJI

KWARTALNIK
ELEKTRONIKI I TELEKOMUNIKACJI

ELECTRONICS AND
TELECOMMUNICATIONS
QUARTERLY

TOM 52 — ZESZYT 3

WARSZAWA 2006

RADA REDAKCYJNA

Przewodniczący

Prof. dr hab. inż. STEFAN HAHN
czł. rzecz. PAN

Członkowie

prof. dr hab. inż. DANIEL JÓZEF BEM — czł. koresp. PAN, prof. dr hab. inż. MICHAŁ BIAŁKO — czł. rzecz. PAN, prof. dr hab. inż. MAREK DOMAŃSKI, prof. dr hab. inż. ANDRZEJ HAŁAŚ, prof. dr hab. inż. JÓZEF MODELSKI, prof. dr inż. JERZY OSIOWSKI, prof. dr hab. inż. EDWARD SĘDEK, prof. dr hab. inż. MICHAŁ TADEUSIEWICZ, prof. dr hab. inż. WIESŁAW WOLIŃSKI — czł. koresp. PAN, prof. dr inż. MARIAN ZIENTALSKI

REDAKCJA

Redaktor Naczelny

prof. dr hab. inż. WIESŁAW WOLIŃSKI

Zastępca Redaktora Naczelnego

doc. dr inż. KRYSTYN PLEWKO

Sekretarz Odpowiedzialny

mgr ELŻBIETA SZCZEPANIAK

ADRES REDAKCJI

00-665 Warszawa, ul. Nowowiejska 15/19 Politechnika, pok. 470
Instytut Telekomunikacji, Gmach im. prof. JANUSZA GROSZKOWSKIEGO

Dyżury Redakcji: środa i piątki, godz. 14–16
tel. (022) 234 77 37

Telefony domowe: Redaktora Naczelnego: 812 17 65

Zast. Red. Naczelnego: 826 83 41

Sekretarza Odpowiedzialnego: 251 02 01

Ark. wyd. 15,0	Ark. druk. 12,75	Podpisano do druku w październiku 2006 r.
Papier offset. kl. III 80 g. B-1		Druk ukończono w październiku 2006 r.

Skład, druk i oprawa: Warszawska Drukarnia Naukowa PAN
00-656 Warszawa, ul. Śniadeckich 8
Tel./fax: 628-87-77

WAŻNY KOMUNIKAT DLA AUTORÓW

Rada Redakcyjna na swoim posiedzeniu w dniu 18 stycznia 2006 r. zaleciła redakcji Kwartalnika Elektroniki i Telekomunikacji zrealizowanie następujących ustaleń:

1. PUBLIKOWANIE NADSYŁANYCH ARTYKUŁÓW WYŁĄCZNIE W JĘZYKU ANGIELSKIM

Począwszy od zeszytu 1'2007 wszystkie artykuły będą publikowane w Kwartalniku w języku angielskim, a tym samym artykuły nadsyłane już obecnie do redakcji muszą być przygotowane przez autorów w tym języku. Artykuły już nadesłane do redakcji i wstępnie zakwalifikowane do publikacji w roku 2007 będą zwrócone autorom w celu opracowania ich w wersji angielskiej.

Każdy artykuł przygotowany w języku angielskim powinien być uzupełniony obszernym (np. dwustronicowym) streszczeniem w języku polskim z przytoczeniem najistotniejszych wzorów oraz powołaniem się na ważne tabele i rysunki. Ponadto należy do tekstu angielskiego dołączyć polskie odpowiedniki tytułów tabel i podpisów pod rysunkami. Artykuły te będą recenzowane, a także weryfikowane pod względem językowym.

2. POKRYWANIE KOSZTÓW WYDAWNICZYCH PUBLIKACJI PRZEZ AUTORÓW

Wprowadza się, począwszy od zeszytu 1'2007, zasadę odpłatności za publikowanie artykułu w Kwartalniku, przy założeniu działalności redakcji non profit, to jest pokrywania przez autorów lub instytucje ich zatrudniające, kosztów wydawniczych w podstawowej wysokości 760 zł za arkusz wydawniczy. Uzyskana w ten sposób kwota będzie przeznaczana na uzupełnienie ograniczonych środków otrzymywanych z PAN na działalność wydawniczą, a w szczególności na zwiększenie objętości kolejnych zeszytów Kwartalnika oraz na weryfikację językową publikowanych w języku angielskim artykułów. Zwiększenie objętości poszczególnych zeszytów jest konieczne ze względu na dużą liczbę nadsyłanych artykułów, których publikowanie w obecnej objętości zeszytów, w sposób nie akceptowalny wydłuża czas ich ukazania się w druku.

W przypadku, wyrażonej w formie pisemnej, prośby autorów o przyspieszenie ich publikacji w stosunku do kolejności wynikającej z terminu ich nadesłania, opłata za publikację artykułu wynosić będzie 1500 zł za arkusz wydawniczy.

W uzasadnionych przypadkach, przedstawionych pisemnie, redakcja może podjąć decyzję częściowego lub całkowitego zwolnienia autorów z podstawowej odpłatności za publikację. Redakcja Kwartalnika wyraża nadzieję, że autorzy artykułów przeznaczonych do opublikowania w roku 2006, podejmą starania wniesienia powyższych opłat już w roku bieżącym. Dotyczy to w szczególności opracowań realizowanych w ramach „grantów”. Wpłaty należy dokonywać przelewem bankowym na konto Warszawskiej Drukarni Naukowej PAN, BPH PBK S.A. w Krakowie VIII O/Warszawa Nr 1060 0076 0000 4010 3000 0 510 z dopiskiem dla Kwartalnika Elektroniki i Telekomunikacji.

REDAKCJA

„Kw
Quarter
Elektro

Kw
demii I
jest cza
prezent
glądow
ktronik
medyc

Au
czeniu

Ar
mi bac
danej
matem
Comm

W
krajow
nauko
w ram
stawia

Cz
krajow
nych.

K
ułatw
lub z
w jęz

N
spraw
publi
Reda

A
stron

Szanowni Autorzy

„Kwartalnik Elektroniki i Telekomunikacji” — Electronics and Telecommunications Quarterly jest kontynuatorem tradycji powstałego 52 lata temu kwartalnika pt. „Rozprawy Elektrotechniczne”.

Kwartalnik jest czasopismem Komitetu Elektroniki i Telekomunikacji Polskiej Akademii Nauk. Wydawany jest przez Warszawską Drukarnię Naukową PAN. Kwartalnik jest czasopismem naukowym, na którego łamach są publikowane artykuły i komunikaty prezentujące wyniki oryginalnych prac teoretycznych i doświadczalnych, a także przeglądowych. Związane są one z szeroko rozumianymi dziedzinami współczesnej elektroniki, telekomunikacji, mikroelektroniki, opoelektroniki, radiotechniki i elektroniki medycznej.

Autorami publikacji są wybitni naukowcy, znani specjaliści o wieloletnim doświadczeniu, a także młodzi badacze — głównie doktoranci.

Artykuły charakteryzują się oryginalnym ujęciem zagadnienia, interesującymi wynikami badań, krytyczną oceną teorii lub metod, omówieniem aktualnego stanu, lub postępu danej gałęzi techniki oraz omówieniem perspektyw rozwojowych. Sposób pisania matematycznej części artykułów zgodny jest z wytycznymi IEC (International Electronics Commission) oraz ISO (International Organization of Standardization).

Wszystkie publikowane w Kwartalniku artykuły są recenzowane przez znanych krajowych specjalistów, co zapewnia że publikacje te są uznawane jako autorski dorobek naukowy. Opublikowane w kwartalniku wyniki prac naukowych zrealizowanych w ramach „GRANTów” Komitetu Badań Naukowych spełnia więc jeden z wymogów stawianych tym pracom.

Czasopismo dociera do wszystkich zajmujących się elektroniką i telekomunikacją krajowych ośrodków naukowych oraz technicznych, a także szeregu instytucji zagranicznych. Jest ponadto prenumerowane przez liczne grono specjalistów i biblioteki.

Każdy Autor otrzymuje bezpłatnie 20 egzemplarzy nadbitek swojego artykułu, co ułatwia przesłanie go do indywidualnych wybranych przez Autora osób i instytucji w kraju lub za granicą. Ułatwia to dodatkowo fakt, że w Kwartalniku są publikowane artykuły w języku angielskim.

Nadesłane do redakcji artykuły są publikowane w terminie około pół roku, w przypadku sprawnej współpracy Autora z Redakcją. Wytyczne dla Autorów dotyczące formy publikacji są zamieszczone w zeszytach Kwartalnika, można je także otrzymać w siedzibie Redakcji.

Artykuły można dostarczać osobiście, lub pocztą pod adresem zamieszczanym na stronie redakcyjnej w każdym zeszycie.

Redakcja

L. Tom
fu
Ł. Śliw
tra
J. Siuz
V.
K. Gór
ko
J. Stan
na
K. Pac
de
K. Gór
w
K. Poż
R. S. F
Z. Nag
n
Z. Nag
Inform

L. Tom
Ł. Śliw
l
J. Siuz
K. Gór
J. Stan
K. Pa
s
K. Gór
c
K. Po
R. S.
Z. Na
Z. Na
Inform

SPIS TREŚCI

L. Tomawski, M. Górnicki: Równoległy obwód rezonansowy o stałym współczynniku dobroci funkcji częstotliwości	299
L. Śliwczyński: Analiza błędzenia progu komparacji oraz bitowej stopy błędów w łączach transmisyjnych z odcięciem składowej stałej	309
J. Siuzdak, T. Czarnecki: Kwalifikacja linii abonenckich do usług ADSL za pomocą modemów V.34	323
K. Górecki, J. Zarębski: Estymacja parametrów modelu termicznego elementów półprzewodnikowych	347
J. Stanclik: Modelowanie właściwości mikromocowych wzmacniaczy operacyjnych w obszarze nasycenia	361
K. Pacholski: Ocena przydatności przetworników pomiarowych wielkości elektrycznych do pomiaru sygnałów odształconych	373
K. Górecki: Wyznaczanie nieizotermicznych charakterystyk dławikowych przetwornic dc-dc o sterowaniu PWM w stanie ustalonym przy wykorzystaniu metody modeli uśrednionych	393
K. Poźniak: Parametryzowany interfejs komunikacyjny dla układów FPGA	411
R. S. Romaniuk, J. Dorosz: Wytwarzanie i charakteryzacja światłowodów kapilarnych	427
Z. Nagórny, A. Kos: Projektowanie topografii systemów VLSI. Cz. 1. Style i etapy projektowania, rozmieszczanie modułów	451
Z. Nagórny, A. Kos: Projektowanie topografii systemów VLSI. Cz. 2. Algorytm min-cut	469
Informacje dla Autorów (w jęz. polskim)	489

CONTENTS

L. Tomawski, M. Górnicki: A parallel resonant circuit with constant quality factor vs. frequency ...	299
L. Śliwczyński: Analysis of baseline wander and BER performance in transmission links with low-frequency removal	309
J. Siuzdak, T. Czarnecki: Subscriber line qualification for ADSL by means of V.34 modems	323
K. Górecki, J. Zarębski: Parameters estimation of thermal model of semiconductor devices	347
J. Stanclik: Modeling of micropower operational amplifiers' proprieties in saturation region	361
K. Pacholski: Estimating of suitability of electric quantity transducers to measurement of deformed signals	373
K. Górecki: Calculations of the nonisothermal characteristics of dc-dc choopers with PWM control at steady-state using average models method	393
K. Poźniak: Parametrized communication interface for FPGA devices	411
R. S. Romaniuk, J. Dorosz: Fabrication and characterization of capillary optical fibres	427
Z. Nagórny, A. Kos: The sedign of the VLSI circuit layout. Part 1. Styles, phases, placement	451
Z. Nagórny, A. Kos: The design of the VLSI circuit layout. Part 2. Min-cut algorithm	469
Information for the Authors (w jęz. ang.)	492

and
in L
FD
for
cap
R_r.

A parallel resonant circuit with constant quality factor vs. frequency

LECH TOMAWSKI, MARCIN GÓRNICKI

*Instytut Fizyki Uniwersytetu Śląskiego w Katowicach
ul. Uniwersytecka 4, 40-007 Katowice
tomawski@us.edu.pl*

*Otrzymano 2005.07.15
Autoryzowano 2005.10.06*

If in a parallel RLC resonant circuit both capacitance and inductance are tuned in such a way that their product changes but their quotient is constant then the resonant frequency $f_r = \frac{1}{2\pi\sqrt{LC}}$ changes, but the quality factor of the circuit $Q = R\sqrt{\frac{C}{L}}$ is constant and independent of the frequency. This interesting case is not used in classical RLC circuits because the capacitor is tuned differently than the inductor but in the resonant circuits with FDNC¹ or FDNR² and similar ones it is easy to find a way to suitably tune both the resonating elements.

Keywords: Bruton's transformation, FDNC, parallel resonant circuit

1. BASIC FORMULAS

Let us consider the parallel resonant circuit formed from FDNC Bruton [1] circuit and resistor R_r according to Fig. 1. Equivalent diagram of the resonant circuit shown in Fig. 1b contains resistor R_r , ideal FDNC element and parasitic elements of the real FDNC circuit the capacitance C_p and the resistance R_p . These parasitic elements are formed by nonideal parameters of the operational amplifiers and losses of the used capacitors. Resistance R_p has been neglected since its value is a few order greater than R_r .

¹ FDNC: Frequency Dependent Negative Conductance: $Y(s) = s^2D$, where D is constant

² FDNR: Frequency Dependent Negative Resistance: $Z(s) = s^2E$, where E is constant.

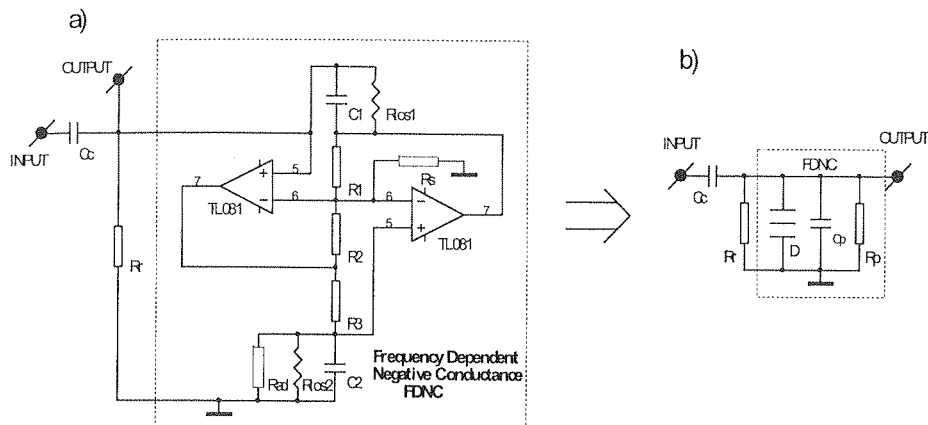


Fig. 1. The scheme of the resonant circuit a), the equivalent diagram of the resonant circuit from Fig. 1a-b)

Rys. 1. Schemat obwodu rezonansowego a), schemat zastępczy obwodu rezonansowego z rys. 1a-b)

Resonant frequency of this circuit is:

$$f_r = \frac{1}{2\pi \sqrt{R_r D}} \quad (1)$$

where D is constant

$$D = \frac{C_1 C_2 R_1 R_3}{R_2} \left[\frac{A s^2}{V} \right] \quad (2)$$

The quality factor of the circuit contained from resistance R_r and FDNC is [2]:

$$Q = \frac{\omega_r D}{C_p} = \frac{1}{C_p} \sqrt{\frac{D}{R_r}} \quad (3)$$

To achieve a constant value of Q both resonating elements R_r and FDNC must be simultaneously tuned. Let us assume that in order to achieve it the resistance R_r equals resistance R_3 and both called R_t will be simultaneously changed. Then formulas (1) and (3) can be rewritten as:

$$f_r = \frac{1}{2\pi R_t \sqrt{\frac{C_1 C_2 R_1}{R_2}}} \quad (4)$$

$$Q = \frac{\sqrt{\frac{C_1 C_2 R_1}{R_2}}}{C_p} \quad (5)$$

Since in formula (5) neither R_r nor R_3 exist the Q value is independent of the frequency tuning. By simultaneous changing of both resistances R_r and R_3 in such a way that $R_r = R_3 = R_t$ the resonant frequency according to formula (4) is inversely proportional to resistance R_t .

Assuming that the AC signal will be fed to the resonant circuit by capacitor C_c , shown in Fig. 1 the equivalent quality factor Q should be described as:

$$Q = \frac{\sqrt{\frac{C_1 C_2 R_1}{R_2}}}{C_p + C_c} \quad (6)$$

2. CAPACITANCE AT THE INPUT OF THE FDNC AND ITS INFLUENCE ON THE QUALITY FACTOR Q

Capacitance C_p shown in Fig.1b exists in the denominator of formulas (3,5,6) and it significantly influences the quality factor Q . The capacitance C_p is the sum of capacitance C_s introduced by losses of capacitors C_1 and C_2 and of the capacitance C_{opa} introduced by the product of gain and bandwidth of the operational amplifiers.

There are a some principles which should be taken into account when calculating all the capacitances present at the input terminals of the circuit.

2.1. ONLY RESISTOR R_R IS USED TO TUNE THE CIRCUIT

In this case the C_s capacitance increases according to formula (7) when the frequency increases:

$$C_s = 2\pi f_r D(\tan \delta_1 + \tan \delta_2) \quad (7)$$

where $\tan \delta_1$ and $\tan \delta_2$ denote dielectric loss factor of the capacitors C_1 and C_2 . In computer simulations these losses were taken into account in the simulation process by connecting resistors R_{los1} and R_{los2} to capacitors C_1 and C_2 in a parallel way. The values of the resistors R_{los1} and R_{los2} were calculated for each frequency from the following formulas: $R_{los1} = 1/\omega_r C_1 \tan \delta_1$ and $R_{los2} = 1/\omega_r C_2 \tan \delta_2$. It was assumed that $\tan \delta_1 = \tan \delta_2 = 0.02$.

The capacitance C_{opa} has a nearly constant value (when $R_1 = R_2$) with frequency and rapidly falls to the region with negative capacitance value. These dependencies are shown in Fig. 2, where C_s , C_{opa} and their sum as capacitance C_p are plotted.

Substituting value C_p into formula (3) the quality factor Q was calculated which is shown in Fig. 3. The transfer characteristics simulated in PSPICE program are also presented in Fig. 3.

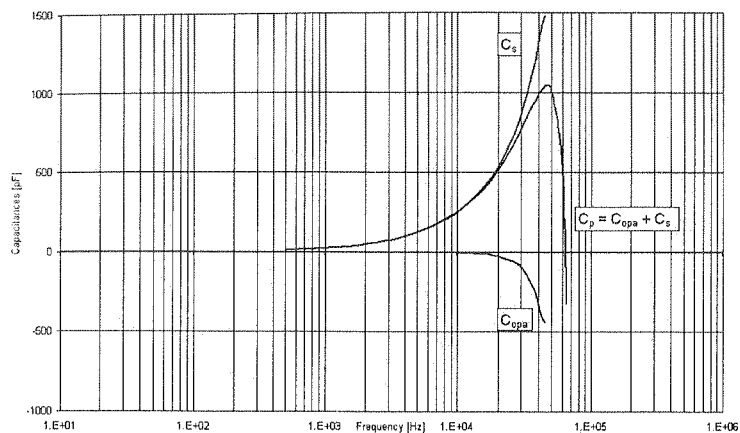


Fig. 2. Capacitances C_s , C_{opa} and their sum C_p vs frequency for the case when only resistor R_r is used for tuning the circuit (from simulations)

Rys. 2. Pojemności C_s , C_{opa} oraz ich suma C_p w funkcji częstotliwości dla przypadku gdy tylko jeden rezystor R_r używany jest do przestrajania obwodu (z symulacji)

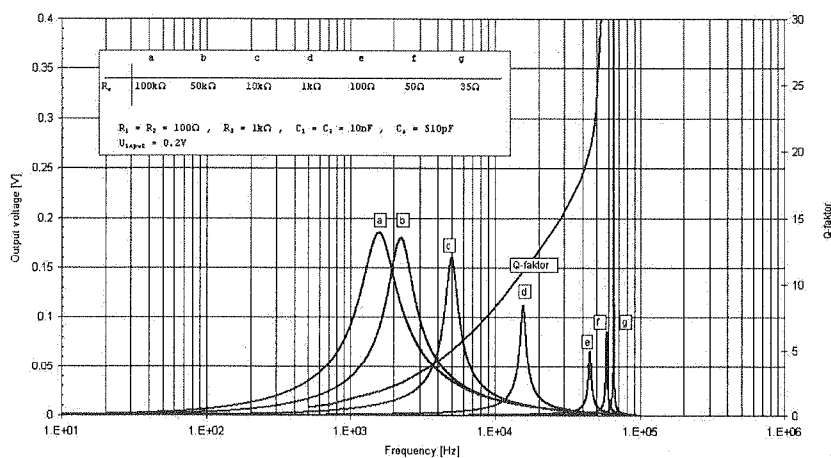


Fig. 3. Q-factor dependence and transfer characteristics for the case when only resistor R_r is used for tuning the circuit. (from simulations)

Rys. 3. Zależność współczynnika dobroci Q oraz krzywe rezonansowe dla przypadku gdy tylko rezystor R_r używany jest do przestrajania obwodu (z symulacji)

2.2. TUNING THE CIRCUIT WITH BOTH R_R AND R_3 RESISTORS

In this case with the frequency increasing the coefficient D decreases and hence the C_s capacitance has a constant value according to formula (7). The C_{opa} capacitance also has a constant value, however at relatively high frequency it falls as well. In Fig. 4 dependencies for C_s , C_{opa} and for sum $C_s + C_{opa} = C_p$ are shown.

The quality factor Q calculated from formula (6) and the simulated transfer characteristics for other values of R_t are shown in Fig. 5. From frequency 25 kHz upwards the quality factor increases, due to the fact that capacitance C_{opa} drops.

The range of the tuning is greater than in point 2.1. Quotient of the resonant frequencies from point 2.1 and 2.2 can be described as: $\frac{R_r}{\sqrt{R_r R_3}}$

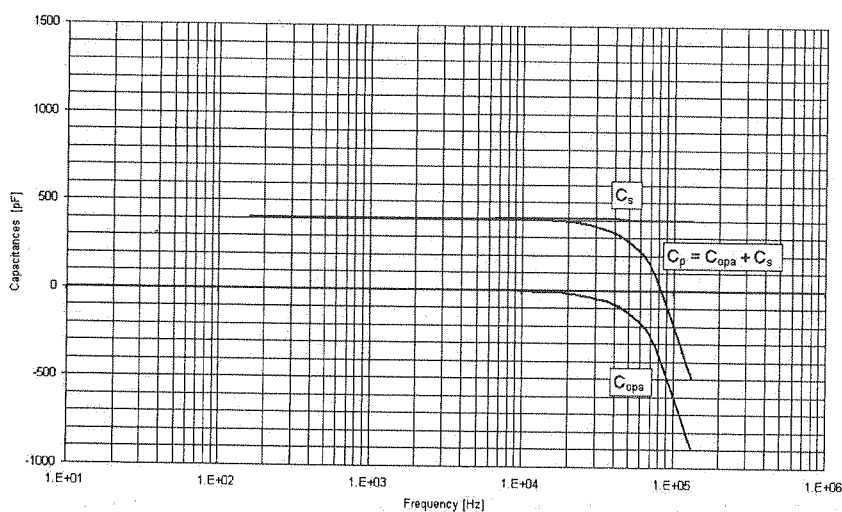


Fig. 4. Capacitances C_s , C_{opa} and $C_s + C_{opa} = C_p$ vs frequency for the case when both resistors R_r and R_3 are simultaneously changed for the tuning of the circuit. (from simulations)

Rys. 4. Pojemności C_s , C_{opa} oraz ich suma $C_s + C_{opa} = C_p$ w funkcji częstotliwości dla przypadku gdy dwa rezystory R_r oraz R_3 są jednocześnie używane do przestrajania obwodu (z symulacji)

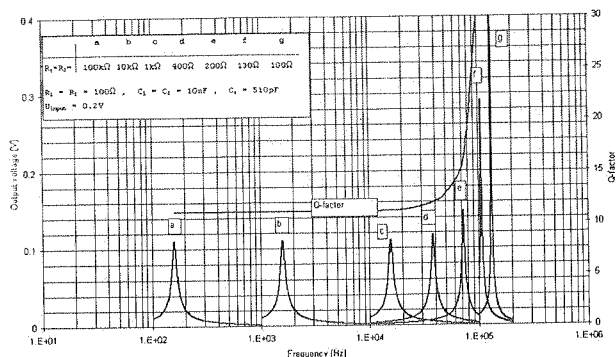


Fig. 5. Q-factor and transfer characteristics vs frequency for the case when both resistors R_1 and R_3 are simultaneously changed for the tuning of the circuit (from simulations)

Rys. 5. Zależność współczynnika dobroci Q oraz krzywe rezonansowe dla przypadku gdy dwa rezystory R_1 oraz R_3 są jednocześnie używane do przestrajania obwodu (z symulacji)

2.3. TUNING THE CIRCUIT WITH BOTH R_R AND R_3 RESISTORS — WITH LOSSES OF ONE OF THE CAPACITORS C_1 OR C_2 INTENTIONALLY INCREASED AT HIGH FREQUENCY REGION

The range in which the quality factor from Fig. 5 remains constant can be extended to the high frequency region if the losses of one (or two) of the capacitors C_1 or C_2 are increased by parallel connecting the additional resistor R_{ad} (shown in Fig. 1). This resistor introduces the additional capacitance C_{ad} whose value can be calculated from formula (7) as:

$$C_{ad} = \frac{C_1 R_1 R_3}{R_2 R_{ad}} \quad (8)$$

where R_{ad} is the resistor connected in a parallel way to the capacitor C_2 . Since in formula (8) there exists resistor R_3 , which influences the frequency, therefore resistor R_{ad} must be changed together with the change of frequency.

2.4. TUNING THE CIRCUIT WITH BOTH R_R AND R_3 RESISTORS — WITH THE NEGATIVE CAPACITANCE REALIZED BY THE SENANI [3] CONCEPT

According to the Senani proposition the adding of resistor R_s , which is shown in Fig. 1, results in introducing the negative capacitance at the input terminals of the circuit whose value can be calculated from the following formula:

$$C_n = -\frac{C_1 R_2}{R_s} \quad (9)$$

2.5. THE COMPLETE FORMULA FOR THE QUALITY FACTOR

After adding capacitances C_{ad} and C_n to the denominator of the formula (6) the complete formula for the quality factor is as follows:

$$Q = \frac{\sqrt{\frac{C_1 C_2 R_1}{R_2}}}{C_C + C_p + C_{ad} + C_n} \quad (10)$$

Sum of capacitances $C_p + C_{ad} + C_n$ can be matched to equal zero in a wide frequency range because C_{ad} is always positive, C_n is always negative and C_p is positive at low frequencies and is negative at high frequencies.

When this sum is equal to zero the quality factor of the basic resonant circuit without the capacitor C_c tends to infinity but if this circuit is fed from the source by the C_c capacitance and if $R_1 = R_2$, and $C_1 = C_2 = C_i$ the quality factor $Q = C_i/C_c$. In our case $C_i = 10\text{nF}$ and $C_c = 510\text{pF}$ and hence $Q \approx 20$. Then output voltage is equal to the input voltage.

3. SIMULATED AND MEASURED RESULTS

Finally simulated transfer characteristics for the circuit are shown in Fig. 6. Output voltage is the same as input voltage (0.2V) from several hundred Hz to nearly 100kHz with TL084 operational amplifiers. Dielectric loss factors of capacitors C_1 and C_2 were

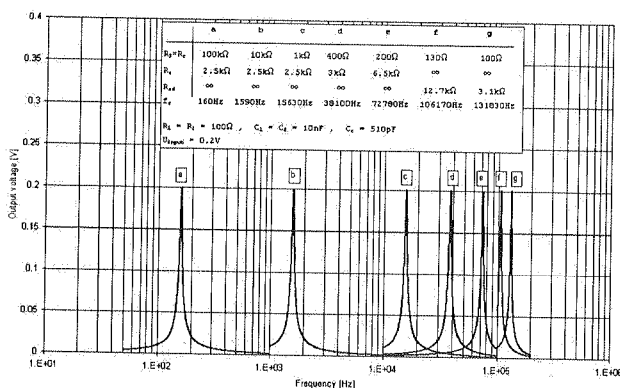


Fig. 6. Transfer characteristics vs frequency for the case when both resistors R_r and R_3 are simultaneously changed for tuning of the circuit and when the resistors R_{ad} and R_s were applied (from simulations)

Rys. 6. Krzywe rezonansowe w funkcji częstotliwości dla przypadku gdy dwa rezystory R_r oraz R_3 są jednocześnie używane do przestrajania obwodu z dodanymi rezystorami R_{ad} oraz R_s (z symulacji)

taken to be $\tan \delta = 0.02$ for both capacitors. It was also assumed that their values do not change with frequency.

Measured dependencies for the same data are shown in Fig. 6. When the coupled capacitance C_c is 100pF (then $Q=100$) the circuit can still be used.

In Fig. 6 and 7 are shown transfer characteristics with a constant voltage amplitude. If resistors R_s and R_{ad} were matched in a slightly different way these transfer characteristics could be set accurately with the constant Q -factor value. Their average Q -factor value is about 20 now.

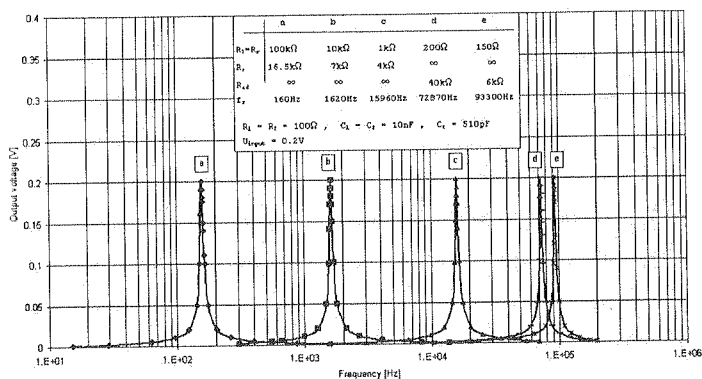


Fig. 7. Some measured transfer characteristics of the circuit

Rys. 7. Zmierzone krzywe rezonansowe

4. RESUME

1. The tuning of the circuit only by resistance R_r is not effective. The range of the tuning is narrow and the value of Q -factor is strongly changed with frequency. There are two reasons of it: Q -factor is dependent on frequency (see formula (3) and Fig.3) and capacitance C_s rises with frequency (see formula (7) and Fig. 2).
2. The use of both resistors R_r and R_3 for tuning the circuit gives a nearly constant value (apart from the high frequency region) of Q -factor (see Fig.5) because according to (6) Q -factor is independent of frequency and the capacitance C_s has a constant value (see Fig.4). The lowest resonant frequency when both resistors R_r and R_3 are used for tuning is about 10 times lower than in the case when only R_r resistor is used.
3. Additional resistor R_{ad} results in decreasing the Q -factor but an additional R_s resistor increases it. In this way the quality factor dependence or amplitude of the transfer characteristics can remain at a constant value.

4. Only three (eventually four) resistors are used to change the resonant frequency of the circuit and all of them can be substituted by digital resistors controlled by digital signals for example from computer PC or microcomputer.

5. REFERENCES

1. L. T. B r u t o n: *Network transfer functions using the concept of Frequency — Dependent Negative Resistance*, IEEE Trans. Circuit Theory, August 1969, pp. 406-408.
2. L. T o m a w s k i, M. M a ł k a: *Resonance phenomena in chosen II and III order circuits*, Electronic and Telecommunications Quarterly, 46 vol. 3, 2000, pp. 339-353.
3. R. S e n a n i: *Alternative modification of the classical GIC structure*, Electron. Lett. No 15, 1996, vol. 32, p. 1329.

L. TOMAWSKI, M. GÓRNICKI

Streszczenie

W wielu praktycznych zastosowaniach obwodów rezonansowych jest istotne utrzymywanie stałej wartości współczynnika dobroci obwodu w jak najszerszym przedziale częstotliwości. Można to osiągnąć przestrajając obydwa rezonujące ze sobą elementy obwodu w taki sposób, aby iloraz ich immitancji pozostawał niezmienny podczas gdy ich iloczyn zmienia się i służy do zmiany częstotliwości. W artykule przebadano równoległy obwód rezonansowy, złożony z rezystancji R_r i elementu FDNC D , zrealizowanego elektronicznie. Ponieważ współczynnik D układu FDNC zależy od 3 rezystorów łatwo jest przestrajac obwód zmieniając odpowiednio wartość jednego z nich oraz wartość rezystora R_r . Układ elektroniczny realizujący FDNC nie jest idealny i dlatego oprócz współczynnika D występuje na zaciskach wejściowych układu pojemność pasożytnicza C_p zależna od strat użytych kondensatorów i pola wzmocnienia wzmacniaczy operacyjnych. Pojemność ta wpływa na obniżenie wartości współczynnika dobroci, a jeżeli stanie się ujemna w obwodzie wzbudzą się drgania niegasnące. Dlatego przebadano zmniejszanie wypadkowej pojemności przez wytwarzanie pojemności ujemnej wg koncepcji Senaniego [3], co jest istotne w przedziale częstotliwości niskich, a także zapobieganie wzbudzeniu przez zwiększanie strat jednego z kondensatorów, co jest ważne w przedziale częstotliwości wysokich. W rezultacie uzyskano bezindukcyjny obwód rezonansowy przestrajany w przedziale od setek Hz do 100 kHz przy stałej wartości dobroci.

Słowa kluczowe: Transformacja Brutona, FDNC (ujemna konduktancja zależna od częstotliwości), równoległy obwód rezonansowy

in
an
th
no
H
ba

Analysis of baseline wander and BER performance in transmission links with low — frequency removal

ŁUKASZ ŚLIWCZYŃSKI

*Institute of Electronics, AGH University of Science and Technology,
30 Mickiewicza Ave., 30-059, Cracow, Poland
e-mail: sliwczyn@agh.edu.pl
phone +48 12 617-27-40
fax +48 12 633-23-98*

Otrzymano 2005.03.16

Autoryzowano 2005.09.30

The paper presents the study of the statistical properties of baseline wander appearing when transmitting digital data through transmission links with low frequency removal. The model of the link is considered where received signal is compared with the threshold obtained by low-pass filtering the non-return to zero digital data signal. Probability density of the signal at the output of different order low-pass filters is investigated. Beta distribution is proposed as the approximation of true probability density function of such signal. Next, the BER performance of transmission link with baseline wander and additive Gaussian noise is analysed. Some practical remarks concerning the localisation of the low frequency poles are also given.

Keywords: baseline wander; BER estimation; sensitivity penalty; random geometric series; digital transmission links

1. INTRODUCTION

In a digital transmission system the data signal supplied to the receiver is processed in analogue fashion up to the point where the detection is made. During detection, the analogue signal is compared against some threshold and the decision is undertaken if the current signal value represents "one" or "zero". Because of presence of random noise, the decision is sometimes incorrect, thus errors are generated during detection. However, not only the noise contributes to the system bit error rate (BER). Limited bandwidth of transmission channel, receiver or transmitter may influence BER through

the phenomena of intersymbol interference (ISI) or baseline wander. The problem of ISI is quite well studied and many useful BER estimations may be found in the literature [1, 2, 3]. This paper is devoted to the study of the baseline wander statistics, its impact on transmission system BER performance and resulting sensitivity penalty.

The origin of the baseline wander lies in the low frequency (LF) cut-off, which is present in almost all continuous-type digital transmission systems. For example the frequency characteristics of analogue part of any fibre optic transmission system has LF cut-off resulting either from inherent features of many laser driving circuits [4] or from exploiting AC coupling somewhere in the signal path. AC coupling of the receiver front-end with the decision circuit is often regarded as a convenient way of generating the threshold level assuming that the probability of „zero” and „one” symbols is equal. But strictly speaking voltage across the coupling capacitor is not constant in this case and depends on the particular pattern of transmitted digital data. This voltage undergoes changes in time what may be regarded as a change of the threshold level. Similar situation exists in links exploiting transformer coupling, e.g. for galvanic isolation purposes.

2. LINK MODEL

The study of the baseline wander in a digital baseband transmission system may be performed considering the equivalent schematic diagram of the analogue part of the link, presented in Fig. 1a. In this figure $H_{\text{tran}}(s)$, $H_{\text{rec}}(s)$ and $H_{\text{tgen}}(s)$ describe LF transfer functions associated with the transmitter, receiver front-end and threshold generating circuitry respectively; n represents the additive Gaussian noise present at the output of the input amplifier. In the considered frequency range the transmission channel only attenuates and delays the data signal. Without the loss of generality, these effects may be neglected here, thus it will be assumed further that the transmission channel transfer function is equal to 1. Schematic diagram presented in Fig. 1 is quite general and not all blocks need to be present in practice. For example, if the receiver front-end is AC-coupled with the decision circuit the threshold level may be assumed to be zero, thus no special threshold generating filter is required in this case. When the receiver front-end is DC-coupled with the decision circuit, the threshold signal is usually generated by simple low-pass filter (LPF).

For the purpose of further analysis, the schematic from Fig. 1a will be transformed into more convenient form. First, we define the equivalent LF link transfer function $H_T(s)$, combining all LF limiting factors present in the link. This idea is shown in Fig. 1b, where the detection process is presented as comparison of the signal d with the threshold u_T , obtained from the data by means of filtration. The form of $H_T(s)$ may be found assuming that the schematics from Fig. 1a and b should be equivalent for differential signal u_D . Neglecting for a while the additive noise source n in Fig. 1a we may write:

$$u_D = d \cdot H_{\text{tran}}(s)H_{\text{rec}}(s)[1 - H_{\text{tgen}}(s)]. \quad (1)$$

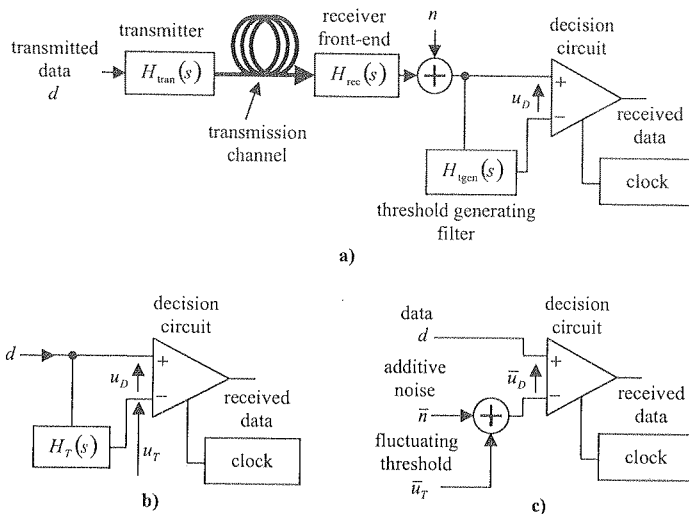


Fig. 1. Equivalent schematic diagram of the analogue part of the baseband link

Rys. 1. Schemat zastępczy analogowej części cyfrowego łącza transmisyjnego

For the differential signal in Fig. 1b we may write:

$$u_D = d \cdot [1 - H_T(s)]. \quad (2)$$

Equating (1) and (2) and rearranging we finally comes to:

$$H_T(s) = 1 - H_{\text{tran}}(s)H_{\text{rec}}(s)[1 - H_{\text{tgen}}(s)]. \quad (3)$$

As it stems from the properties of $H_{\text{tran}}(s)$, $H_{\text{rec}}(s)$ and $H_{\text{tgen}}(s)$ the transfer function $H_T(s)$ is always of the low-pass type. For the link functionality its order must be at least one, but up to three poles may appear in different practical realisations. Assuming non-return to zero (NRZ) signalling with bit duration T_b the equation describing data signal is:

$$d(t) = \sum_{i=-\infty}^{+\infty} a_i \Pi(t - iT_b) \quad (4)$$

where a_i may take value 0 or A with equal probabilities $P = 0.5$ and $\Pi(\cdot)$ denotes the pulse taking value one within the interval $(-T_b/2, T_b/2)$ and equal to zero outside this interval. After passing this signal through the block with transfer function $H_T(s)$ we get the signal (see Fig. 1b):

$$u_T(t) = \int_{-\infty}^{+\infty} h_T(\nu) \sum_{i=-\infty}^{+\infty} a_i \Pi(t - \nu - iT_b) d\nu. \quad (5)$$

where $h_t(t)$ denotes the impulse response related to $H_T(s)$. Rearranging equation (5) we may write:

$$u_T(t) = \sum_{i=-\infty}^{+\infty} a_i r_T(t - iT_b), \quad (6)$$

where:

$$r_T(t) = \int_{-\infty}^{+\infty} h_T(v) \Pi(t - v) dv. \quad (7)$$

Detection in the receiver is performed using the samples of the differential signal $u_D(t)$, taken in the discrete time instants $t_j = jT_b$, where j is some integer number. Because of stationarity of such process [5] it is possible to limit our attention to any particular value of j , say $j = 0$. Thus, exploiting the causality of $h_T(t)$ we may write:

$$u_T[0] = \sum_{i=-\infty}^{-1} a_i r_T[-i] + a_0 r_T[0]. \quad (8)$$

The square brackets are used to distinguish that we are considering here only the samples of continuous-time signals taken with the rate $R_b = 1/T_b$. From equation (8) we may see that some correlation exists between currently detected bit a_0 and the threshold signal $u_T[0]$. To cope with this problem we will write the equation for the differential signal from Fig. 1b, collecting together components dependent on currently detected bit a_0 :

$$u_D[0] = a_0 - u_T[0] = a_0(1 - r_T[0]) - \sum_{i=-\infty}^{-1} a_i r_T[-i]. \quad (9)$$

Writing equation (9) this way we distinguished the component in the differential signal $u_D[0]$ that does not depend on a_0 . It may be said that decision about currently detected bit is undertaken comparing the component $a_0(1 - r_T[0])$ with some threshold value represented by the summation in equation (9). For convenience of further analysis we divide both sides of equation (9) by factor $(1 - r_T[0])$ and additionally exploit the fact that assuming uncorrelatedness of transmitted bits it is unimportant from the statistical viewpoint if we put a_{-i} or a_i . The normalised summation displaying no correlation with currently detected bit will be called the fluctuating threshold and will be defined as:

$$\bar{u}_T = \frac{\sum_{i=1}^{+\infty} a_i r_T[i]}{1 - r_T[0]} = \sum_{i=1}^{+\infty} a_i \bar{r}_T[i], \quad (10)$$

where $\bar{r}_T[i] = r_T[i]/(1 - r_T[0])$.

The complete equivalent model of the link in the LF range is presented in Fig. 1c. In this figure the data signal d is compared with the noise corrupted threshold, composed of the fluctuating threshold $\bar{u}_T$ defined above and the additive Gaussian noise $\bar{n}$. Remembering the normalisation used in equation (10) the additive noise variance may be written as:

$$\sigma_{\bar{n}}^2 = \frac{\int_0^{+\infty} N(f) |1 - H_T(f)|^2 df}{(1 - r_T[0])^2}, \quad (11)$$

where $N(f)$ denotes one-sided spectral density of the noise process n (see Fig. 1a). In the next section some important statistical parameters of the fluctuating threshold $\bar{u}_T$ will be derived.

3. STATISTICAL PROPERTIES OF THE FLUCTUATING THRESHOLD

3.1. MEAN AND VARIANCE

Mean value and variance of the fluctuating threshold $\bar{u}_T$ may be derived basing on its definition given in equation (10). Exploiting the fact shown in [6] that $\sum_{i=1}^{+\infty} \bar{u}_T[i] = 1$, we may derive the mean value of $\bar{u}_T$ as being equal:

$$E\{\bar{u}_T\} = E\{a_n\} = \frac{A}{2} \quad (12)$$

It stems from equation (13) that $E\{\bar{u}_T\}$ is independent on the particular form of the transfer function $H_T(s)$.

The variance of $\bar{u}_T$ may be calculated to be:

$$\sigma_T^2 = \frac{A^2}{4} \sum_{i=1}^{+\infty} \bar{r}_T^2[i] \quad (13)$$

So to calculate σ_T^2 it is enough to know only the samples of $\bar{r}_T$. Because $H_T(s)$ may be expressed as a low-order rational function the analytical expressions may always be written for σ_T^2 . Such expressions are quite complicated, unfortunately. It appears however that σ_T^2 is influenced the most by the part of the circuit with the lowest time constant (or equivalently with the highest LF cut-off) [6]. Thus, the most important cases for practice are those with just one pole or with two or three closely located poles. Assuming that the poles have identical time constants the approximated expressions (with accuracy better than 1%) for fluctuating threshold variance are [6]:

$$\text{single pole} \quad \sigma_T^2 = \frac{A^2}{4} \frac{1 - \beta}{1 + \beta} = \frac{A^2}{4} \operatorname{tgh}\left(\frac{\xi}{2}\right) \approx 0.125 A^2 \xi, \quad (14)$$

$$\text{double pole} \quad \sigma_T^2 \approx A^2 (0.313\xi + 0.125\xi^2 + 0.011\xi^3), \quad (15)$$

$$\text{triple pole} \quad \sigma_T^2 \approx A^2 (0.515\xi + 0.422\xi^2 + 0.245\xi^3). \quad (16)$$

In equations (14)-(16) parameter ξ denotes the normalised bit duration and is defined as:

$$\xi = T_b/\tau_P \quad (17)$$

where τ_P is the time constant associated with the LF poles.

3.2. PROBABILITY DENSITY FUNCTION

Before beginning with the problem of finding the fluctuating threshold PDF we first note some general properties that have to be satisfied by this function, regardless of the number of LF poles present in $H_T(s)$. These properties may be deduced directly from the definition given by equation (10) [6] and will be helpful in further considerations:

— $A/2$ is the symmetry point of $\bar{u}_T$ PDF;

— $\bar{u}_T$ is limited to the range $\left(\frac{A}{2} \left(1 - \sum_{i=1}^{\infty} |\bar{r}_T[i]| \right), \frac{A}{2} \left(1 + \sum_{i=1}^{\infty} |\bar{r}_T[i]| \right) \right)$.

In case of just one LF pole we may write the expression for the samples of $\bar{r}_T[i]$ in the form:

$$\bar{r}_T[i] = (1 - \beta)\beta^{i-1}, \quad (18)$$

where $\beta = \exp(-T_b/\tau_P)$. Thus using equation (10) we may write:

$$\bar{u}_T = \sum_{i=1}^{+\infty} a_i (1 - \beta)\beta^{i-1}. \quad (19)$$

This equation represent what is known in the literature as random geometric series [8, 9]. Substantial amount of work has been done in the past to develop methods of calculating PDF of such series but obtained algorithms are rather complex and solutions in the closed form are limited only to some particular β values. It is possible however to obtain quite simple approximation of $\bar{u}_T$ PDF exploiting the property of semicontraction mapping.

To proceed further we rearrange equation (19) in the following form:

$$\begin{aligned} \bar{u}_T &= (1 - \beta) \sum_{i=1}^{+\infty} a_i \beta^{i-1} = (1 - \beta) a_1 + (1 - \beta) \sum_{i=2}^{+\infty} a_i \beta^{i-1} = \\ &= (1 - \beta) a_1 + \beta (1 - \beta) \sum_{i=1}^{+\infty} a_{i+1} \beta^{i-1} = \\ &= w + \beta z, \end{aligned} \quad (20)$$

where $w = (1 - \beta)a_1$ and $z = (1 - \beta) \sum_{i=1}^{+\infty} a_{i+1} \beta^{i-1}$ are two independent random variables.

We may note that from the statistical viewpoint summations $\sum_{i=1}^{+\infty} a_i \beta^{i-1}$ and $\sum_{i=1}^{+\infty} a_{i+1} \beta^{i-1}$ are identical. Thus PDFs of variables $\bar{u}_t$ and z must be identical.

Equation (20) describes the fluctuating threshold as a sum of two independent random variables. PDF of variable w is binomial and may be written as:

$$f_w(x) = 0.5\delta(x) + 0.5\delta(x - A(1 - \beta)), \quad (21)$$

where $\delta(\cdot)$ stands for Dirac delta. PDF of variable z is identical with PDF of fluctuating threshold itself. Using the fact that PDF of sum of two independent random variables is the convolution of individual PDFs we may derive the following equation that must be obeyed by $\bar{u}_T$ PDF:

$$f_T(x) = \frac{0.5}{\beta} \left[f_T\left(\frac{x}{\beta}\right) + f_T\left(\frac{x - A(1 - \beta)}{\beta}\right) \right], \quad (22)$$

where $f_T(\cdot)$ denotes PDF of $\bar{u}_T$ and we used the fact that $f_z(\cdot) = f_T(\cdot)$. Integrating both sides of equation (22) we may write it using cumulative distribution function (CDF) $F_T(\cdot)$ instead of PDF. Exploiting the symmetry of fluctuating threshold CDF $F_T(x) = 1 - F_T(A - x)$ we finally get:

$$F_T(x) = 0.5 \left[F_T\left(\frac{x}{\beta}\right) + 1 - F_T\left(\frac{A - x}{\beta}\right) \right]. \quad (23)$$

From the properties of $\bar{u}_T$ in the link with one LF pole it stems that $F_T(x) = 0$ for $x < 0$ and $F_T(x) = 1$ for $x > A$.

Equation (23) may be regarded as a type of mathematical operation transforming function $F_T(\cdot)$ onto itself. It may be proved [8, 9] that transformation defined by the right hand side of equation (23), i.e.:

$$T\{F(x)\} = 0.5 \left[F\left(\frac{x}{\beta}\right) + 1 - F\left(\frac{A - x}{\beta}\right) \right] \quad (24)$$

forms the semicontraction mapping. Thus taking any continuous function monotonically increasing from 0 to 1 in the interval $[0, A]$ and putting it into equation (24) the approximation of $\bar{u}_T$ CDF may be found by recursion. It stems from the properties of transformation (24) that after 100 iterations maximum error of approximation will be smaller than 10^{-32} [9].

Results of numerical calculations performed accordingly to equation (24) for normalised bit duration $\xi = 0.05$ are presented in Fig. 2. It appears that so-called beta distribution [10] approximates very well distribution $f_T(\bar{u}_T)$ in the case of first order

link. The similarity of both distributions is so great that they are indistinguishable when plotted using neither linear nor semilogarithmic scales (see Fig.3 a and b). Assuming more generally that $\bar{u}_t$ is nonzero in the interval $[\bar{u}_{T\min}, \bar{u}_{T\max}]$ the beta distribution may be described by the following equation:

$$f_T(x) = \begin{cases} \frac{(x - \bar{u}_{T\min})^{\gamma-1} (x - \bar{u}_{T\max})^{\gamma-1}}{(\bar{u}_{T\max} - \bar{u}_{T\min})^{2\gamma-1} B(\gamma, \gamma)}; & \text{for } x \in [\bar{u}_{T\min}, \bar{u}_{T\max}] \\ 0; & \text{outside this interval} \end{cases} \quad (25)$$

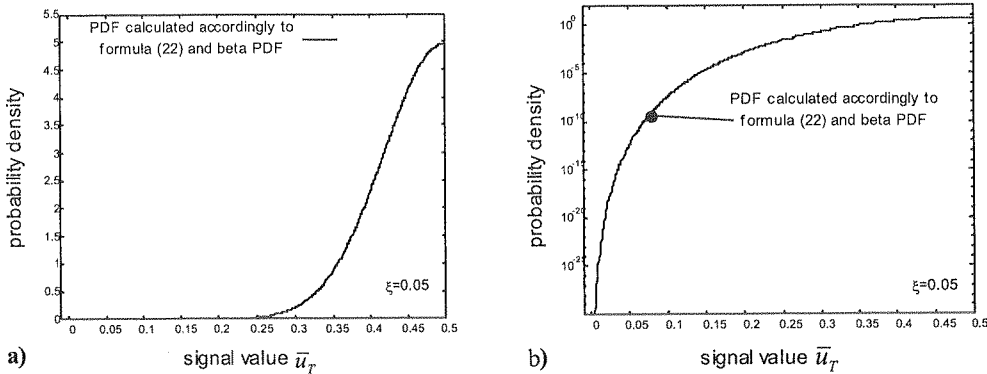


Fig. 2. Example of PDF calculated using semicontraction mapping accordingly to equation (22): plotted using linear scales (a) and using semilogarithmic scales (b)

Rys. 2. Przykładowe funkcje gęstości prawdopodobieństwa wyznaczone za pomocą metody odwzorowania związanego (zgodnie z równaniem (22)): w skali liniowej (a) oraz w skali półlogarytmicznej (b)

where $B(\gamma, \gamma) = \int_0^1 t^{\gamma-1} (1-t)^{\gamma-1} dt$ is Euler's beta function and γ is the parameter of the distribution. It is related to the variance of beta distribution by the equation:

$$\gamma = \frac{1}{2} \left[\frac{(\bar{u}_{T\max} - \bar{u}_{T\min})^2}{4\sigma_T^2} - 1 \right], \quad (26)$$

or in case of the first order link:

$$\gamma = \frac{\beta}{1-\beta}. \quad (27)$$

Above equation stems from equation (14) and from the fact that $\bar{u}_{T\max} - \bar{u}_{T\min} = A$. In the strict sense, beta distribution is only the approximation of true distribution of

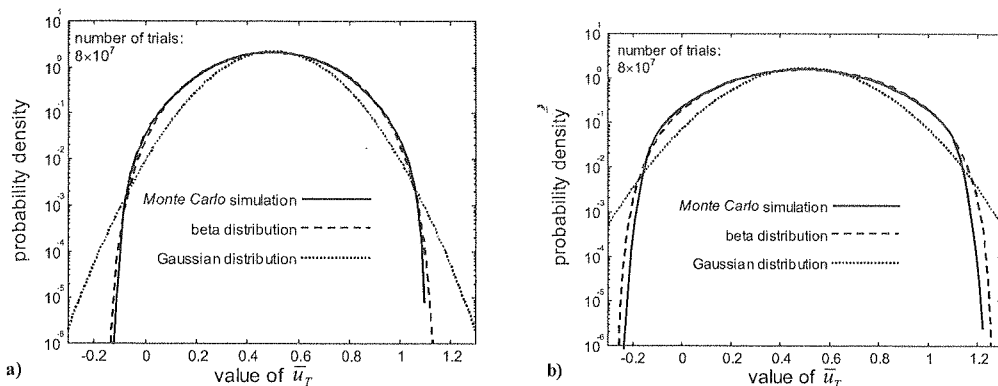


Fig. 3. Results of *Monte Carlo* simulations of fluctuating threshold PDF: two-pole case (a) and three-pole case (b) Parameter $\xi = 0.1$

Rys. 3. Rozkłady prawdopodobieństwa fluktuacji progu komparacji otrzymane metodą *Monte Carlo*: układ z dwoma biegunami (a) i z trzema biegunami (b)

the fluctuating threshold because it does not fulfil equation (22) exactly [6]. It may be proved however that this distribution describes statistical fluctuations of signal at the output of the first order LPF assuming, that the input signal takes value 0 or A and duration of these pulses is exponentially distributed [11]. This is not exact in case of binary data signal as only discrete time duration of the pulses is allowed here, but the probability for the bit to last for time nT_b decreases in the exponential manner. In this sense, statistical behaviour of such two signals is similar.

In case of links with more than one LF pole it is impossible to write the equation similar to (20). However, we postulate here that the approximation of this PDF beta distribution may still be used. The only difference in comparison to the single pole link is that the range of values $\bar{u}_T$ may take is broader than $[0, A]$. The definition of beta distribution used in equation (25) takes this into account.

Theoretically, it is possible to find analytical expressions for extreme values of $\bar{u}_T$ for links with multiple LF poles, but these expressions are very complex. It appears however that simple linear functions of normalised bit duration may be used when poles have the same time constants [6]:

$$\text{double pole} \quad \bar{u}_{T\max} \approx A(0.14\xi + 1.14), \quad (28)$$

$$\text{triple pole} \quad \bar{u}_{T\max} \approx A(0.35\xi + 1.23). \quad (29)$$

Value of $\bar{u}_{T\min}$ may be obtained from the relationship $\bar{u}_{T\min} = A - \bar{u}_{T\max}$. To check how beta distribution performs when approximating fluctuating threshold PDF we used *Monte Carlo* methods to estimate $f_i(\bar{u}_T)$ for links with two and three identical LF poles [6]. Standard *Monte Carlo* method was used to estimate the distribution near

the mean value. Calculations were performed accordingly to equation (10) assuming equal probability of transmitting „one” or „zero” and limiting the number of terms in the summation up to 500. Example result for link with two poles is presented in Fig. 3a and for link with three poles in Fig. 3b along with beta distribution and Gaussian distribution, commonly applied to approximate statistical behaviour of different complex phenomena [12]. It may be noticed that beta distribution matches the results of *Monte Carlo* simulations much better than the Gaussian one. These results were also confirmed by *importance sampling* [13] *Monte Carlo* simulations, allowing more precise estimating of fluctuating threshold PDF in the regions of extreme $\bar{u}_T$ values [6].

4. ESTIMATION OF BIT ERROR RATE

In this section we will exploit approximation of the fluctuating threshold PDF developed above to assess link BER performance in the presence of additive Gaussian noise.

The probability of error in case of noise corrupted threshold may be written as (see Fig. 1c):

$$P_e = 0.5 [\Pr\{0 > \bar{n} + \bar{u}_T\} + \Pr\{A < \bar{n} + \bar{u}_T\}]. \quad (30)$$

First part of equation (30) describes the probability of erroneous detection of „zero” while the second one the probability of erroneous detection of „one”. Exploiting symmetries of noise and fluctuating threshold PDFs the probability of error from equation (30) may be written in the form:

$$P_e = 0.5 \int_{\bar{u}_{T\min}}^{\bar{u}_{T\max}} \text{erf} c \left(\frac{x}{\sigma_{\bar{n}} \sqrt{2}} \right) f_T(x) dx. \quad (31)$$

Assuming for the moment that there is no fluctuations of the threshold (thus $f_T(x) = \delta(x - A/2)$) we may derive standard expression for BER in link with Gaussian noise from equation (31):

$$P_e = 0.5 \text{erf} c \left(\frac{A}{\sigma_{\bar{n}} 2 \sqrt{2}} \right) = 0.5 \text{erf} c \left(\frac{S/N}{2 \sqrt{2}} \right), \quad (32)$$

where $S/N = A/\sigma_n$ denotes the signal to noise ratio at the input of the decision circuit.

Using definition of beta PDF given in (25) equation (31) may be transformed into more convenient form:

$$P_e = [2^{2\gamma} B(\gamma, \gamma)]^{-1} \int_{-1}^1 (1 - x^2)^{\gamma-1} \text{erf} c \left(\frac{S/N}{2 \sqrt{2}} \frac{1 + x/\bar{A}}{2} \right) dx, \quad (33)$$

where:

$$\bar{A} = A/(\bar{u}_{T\max} - \bar{u}_{T\min}) = A/(2\bar{u}_{T\max} - A) \quad (34)$$

denotes the normalised bit amplitude and parameter γ is defined by equation (26). Although equation (33) is not in the closed form but BER may be quite easily calculated with help of any standard mathematical package, like for example Matlab or Mathematica.

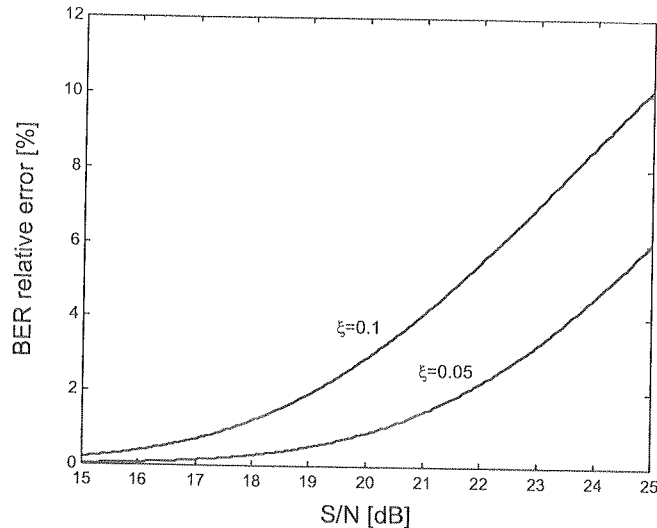


Fig. 4. Relative error of BER resulting from approximating PDF of the fluctuating threshold by beta distribution

Rys. 4. Względny błąd aproksymacji BER przy aproksymacji fluktuacji progu komparacji rozkładem beta

First, to gain some insight into the accuracy of proposed approximations, we computed BER accordingly to equation (31), treating distribution of $\bar{u}_T$ obtained by semicontraction mapping as the exact PDF. The results were compared with that calculated from equation (33) for a few cases. Relative error resulting from this comparison defined as $(BER_{exact} - BER_{beta})/BER_{exact}$ is plotted in Fig. 4. It may be seen that applying beta distribution to approximate fluctuating threshold PDF results in only slight BER underestimation, not exceeding 10% in analysed cases. Accuracy of the approximation increases when normalised bit duration ξ decreases.

Results of BER calculations performed accordingly to equation (33) are presented on the left side of Fig. 5 for links with one (a), two (b) and three (c) identical poles. Normalised bit duration ξ is the parameter for each graph. BER for the link with no baseline wander is plotted with the dashed line. On the right side of Fig. 5 the sensitivity penalty resulting from baseline wander is plotted versus normalised bit duration ξ . Graphs are drawn for BER values equal to 10^{-6} , 10^{-9} and 10^{-12} . From

presented graphs it may be seen that presence of LF poles in the transmission link results in BER degradation. BER increase is more severe when the normalised bit duration is low and when more LF poles are present in the link. In case of link with one pole values of $\xi < 0.01$ may be roughly considered as being acceptable from practical point of view. Such values result in sensitivity penalty lower than about 1.5 dB even for $\text{BER} = 10^{-12}$. For links with two poles similar criterion calls for $\xi < 4 \cdot 10^{-3}$ and for links with three poles condition $\xi < 2 \cdot 10^{-3}$ should be satisfied.

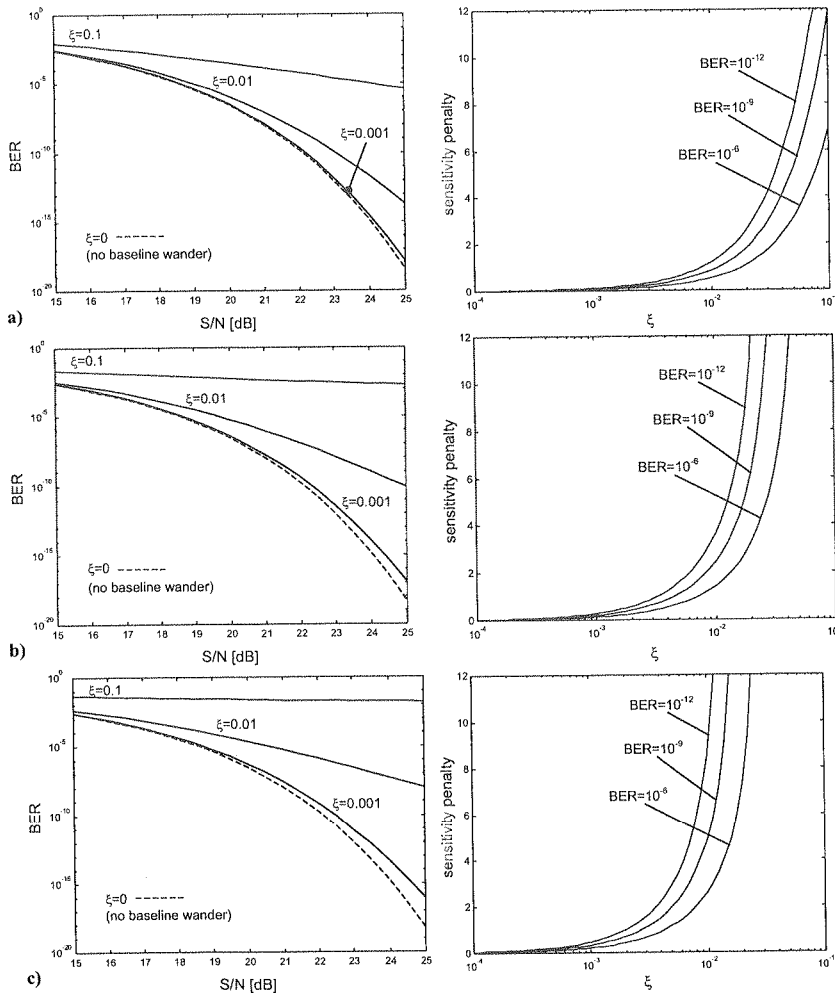


Fig. 5. BER (left) and sensitivity penalty (right) in transmission links with baseline wander: link with one LF pole (a), with two LF poles (b) and with three LF poles (c)

Rys. 5. BER (lewa strona) oraz strata czułości łącza (prawa strona) w łączu transmisyjnym: łącze z pojedynczym biegunem (a), z podwójnym biegunem (b) i z potrójnym biegunem (c)

On the basis of presented graphs the conclusion may also be drawn that baseline wander is almost imperceptible if the time constant of the LF pole is about 300 times greater than the bit period for link with only one pole. In case of two poles similar condition is $\tau_p > 1000T_b$, and for link with three poles $\tau_p > 2000T_b$.

5. CONCLUSIONS

In the paper, a new method of calculating BER in transmission link with baseline wander and additive Gaussian noise for NRZ signalling has been presented. The method is based on approximating the exact solution of functional equation describing baseline wander by beta PDF. Accuracy of this approximation was verified in case of link with one LF pole and resulting BER underestimation appeared not to exceed 10%. This accuracy improves when the amount of baseline wander lowers. For link with more than one pole suitability of beta PDF for modelling fluctuating threshold PDF was verified using statistical *Monte Carlo* methods. The general observation may be made that the beta distribution matches fluctuating threshold PDF much closer than commonly exploited Gaussian distribution.

Proposed approximation allowed calculation of BER degradation in links with baseline wander and resulting sensitivity penalty. Basing on performed calculations some practical rules concerning values of LF time constants in relation to transmission rate were formed.

6. REFERENCES

1. R.D. Gitlin, J.F. Hayes, S.B. Weinstein: *Data communications principles*, Plenum Press, 1992.
2. V.K. Prabhu: *Interference analysis and performance of linear digital communication systems*, chapter 9 in: K. Feher (ed.): *Advanced digital communications*, Prentice-Hall, Englewood Cliffs, 1987.
3. P. Krehlik, Ł. Śliwczyński: *A new approach to BER estimation for ISI corrupted base-band transmission systems*, International Journal of Communication Systems, 2001, no. 14, pp. 513–520.
4. Ł. Śliwczyński, M. Lipiński, P. Krehlik, A. Wolczko: *Performance Analysis of CW Laser Driving Circuits*, Opto-Electronics Review, 1998, vol. 6, no. 2, pp. 131–139.
5. H. Meyr, M. Moeneclaey, S.A. Fechtel: *Digital communication receivers*, New York, Wiley, 1998.
6. Ł. Śliwczyński: *Fibre optic transmission of signals with fluctuating local digital sum*, (in polish), PhD thesis at the AGH University of Science and Technology, Cracow, Poland, 2001.
7. A. Papoulis: *Probability, random variables and stochastic processes*, New York, McGraw-Hill, 1991.
8. F.S. Hill, M.A. Blanco: *Random geometric series and intersymbol interference*, IEEE Trans. Inf. Theory, 1973, vol. 19, no. 3, pp. 326–335.
9. P.J. Smith: *The distribution functions of certain random geometric series concerning intersymbol interference*, IEEE Trans. Inf. Theory, 1991, vol. 37, no. 6, pp. 1657–1662.
10. J.G. Hahn, S.S. Shapiro: *Statistical models in engineering*, New York, Wiley, 1994.

11. R.F. Pawula, S.O. Rice: *On filtered binary processes*, IEEE Trans. on Inf. Theory, 1986, vol. 32, no. 1, pp. 64–72
12. A.M. Street, K. Samaras, D.C. O'Brien, D.J. Edwards: *Closed form expression for baseline wander effects in wireless IR applications*, Electron. Lett., 1997, vol. 33, no 2, pp. 1060–1062
13. D. Lu, K. Yao: *Improved importance sampling technique for efficient simulation of digital communication systems*, IEEE J. on Select. Areas in Commun., 1988, vol. 6, no. 1, pp. 67–75

L. ŚLIWCZYŃSKI

ANALIZA BŁĄDZENIA PROGU KOMPARACJI ORAZ BITOWEJ STOPY BŁĘDÓW W ŁĄCZACH TRANSMISYJNYCH Z ODCIĘCIEM SKŁADOWEJ STAŁEJ

Streszczenie

Artykuł dotyczy analizy zjawiska błędzenia progu komparacji występującego w cyfrowych łączach transmisyjnych z odciętą składową stałą. W przedstawionym modelu zastępczym łącza w zakresie małych częstotliwości odbierany sygnał NRZ jest porównywany z progiem komparacji, otrzymywanym poprzez dolnoprzepustową filtrację sygnału danych. Transmitancja filtru na wyjściu którego otrzymujemy próg komparacji uwzględnia ograniczenia dolnopasmowe, wnoszone do łącza przez układy nadajnika oraz wzmacniacza odbiorczego. Na podstawie przeprowadzonych rozważań teoretycznych zaproponowano aproksymację funkcji gęstości prawdopodobieństwa sygnału na wyjściu filtru w postaci rozkładu beta. Poprawność zaproponowanego rozkładu sprawdzono w przypadku transmitancji I rzędu przy użyciu metody odwzorowania zwężającego. W przypadku łącz o transmitancji II i III rzędu celem sprawdzenia poprawności aproksymacji posłużono się metodami *Monte Carlo* (w wersji standardowej oraz *importance sampling*). Następnie, wykorzystując zaproponowaną aproksymację dokonano analizy wpływu fluktuacji progu komparacji na bitową stopę błędów oraz wyznaczono redukcję czułości łącza transmisyjnego. Na zakończenie zostały podane praktyczne uwagi odnośnie sposobu doboru stałych czasowych związanych z ograniczeniami dolnopasmowymi w łączu transmisyjnym.

Słowa kluczowe: błędzenie progu komparacji, stopa błędów, redukcja czułości łącza, losowe szeregi geometryczne, cyfrowe łącza transmisyjne

Najp
roko
syme
(ang.
mode

Kwalifikacja linii abonenckich do usług ADSL za pomocą modemów V.34

JERZY SIUZDAK, TOMASZ CZARNECKI

*Instytut Telekomunikacji Politechniki Warszawskiej,
Warszawa, Nowowiejska 15/19
e-mail: siuzdak@tele.pw.edu.pl.*

*Otrzymano 2005.08.08
Autoryzowano 2005.10.19*

W pracy przedstawiono metodę kwalifikacji linii abonenckich do usług ADSL w oparciu o rezultaty sondowania linii modemami pracującymi w paśmie głosowym zgodnie z protokołem V.34. Istota tej metody polega na powiązaniu nachylenia charakterystyki przenoszenia toru (zdominowanego przez badaną pętlę abonencką) w paśmie głosowym z jego tłumieniem na częstotliwości 300 kHz. Charakterystyka ta jest mierzona przez modemy V.34 na etapie inicjalizacji połączenia. W pracy uzyskano odpowiednie zależności teoretyczne, które poprzez opracowany system pomiarowy wraz z oprogramowaniem zostały następnie zweryfikowane eksperymentalnie. Uzyskane wyniki pozwalają się spodziewać błędu w określeniu tłumienia linii abonenckiej na częstotliwości 300 kHz na poziomie nie większym niż ± 5 dB. Mając obliczone tłumienie linii na częstotliwości 300 kHz można dokonać oszacowania spodziewanych do uzyskania przepływności w obydwu kierunkach transmisji w systemie ADSL. Zastosowano tu przybliżenia funkcją sigmoidalną z tak dobranymi parametrami, aby uzyskać zadane prawdopodobieństwo tego, iż rzeczywista przepływność linii nie będzie mniejsza od przepływności obliczonej.

Słowa kluczowe transmisja danych, sieci dostępowe, ADSL, V.34, pętla abonencka, kwalifikacja linii, tłumienność

1. WSTĘP

Najpowszechniejszy obecnie trend dostarczania abonentom telefonicznym usług szerokopasmowych polega na wykorzystaniu istniejącego okablowania (kable z parami symetrycznymi) i zastosowaniu szybkich modemów cyfrowych w jednej z technik DSL (ang. DSL-digital subscriber line). Najbardziej obecnie rozwinięta technologia takich modemów to technologia ADSL (ang. Asymmetric DSL) wykorzystująca modulację

DMT (ang. discrete multitone), która pozwala na osiągnięcie przepływności ponad 800 kbit/s w kierunku od abonenta (ang. upstream) oraz do ok. 8 Mbit/s w kierunku do abonenta (ang. downstream). Technologia ta jest już dobrze rozwinięta i została znormalizowana w zaleceniach ITU G.992.1 oraz G.992.2. Największym problemem, z jakim borykają się operatorzy przy wprowadzaniu usług szerokopasmowych (xDSL) jest to, że istniejące okablowanie nie zawsze jest odpowiednie do szybkiej transmisji danych. Warto tu podkreślić, że poprawność pracy telefonii głosowej ($f < 4$ kHz) nie gwarantuje poprawności pracy systemów ADSL ($26 \text{ kHz} < f < 1104 \text{ kHz}$) ze względu na różne zakresy częstotliwości. Dlatego przed instalacją systemu ADSL niezbędnym krokiem jest tzw. kwalifikacja linii polegająca na sprawdzeniu, czy na danej linii abonenckiej system ADSL będzie pracował poprawnie przy założonych przepływnościach. W niniejszej pracy opisano metodę ekstrapolacji danych pomiarowych (umożliwiających kwalifikację linii) z zakresu telefonii głosowej na pasmo ADSL oraz jej implementację za pomocą pary modemów działających w paśmie głosowym. Pozwala to dokonać kwalifikacji linii na podstawie pomiarów przeprowadzanych z obydwu jej końców, jednak bez angażowania personelu utrzymania. Podobna kwalifikacja linii na podstawie pomiarów dwustronnych w paśmie głosowym dokonywanych przez modemy V.34 jest proponowana i realizowana przez amerykańską firmę *Telcordia* (system *Sapphire*) [16]. Należy jednakże podkreślić, iż istnieją istotne różnice między niniejszą metodą a pracą [16]. Dotyczy to zarówno układu pomiarowego w centrali (w bieżącej pracy zastosowano modem analogowy, w pracy [16] — modem cyfrowy), a także samej metody szacowania przepływności. W naszym rozwiązaniu stosuje się zaprezentowane dalej podejście statystyczne, natomiast w pracy [16] szacowania dokonuje się na podstawie bazy danych o odnośnych poziomach zakłóceń. Plan niniejszej pracy jest następujący. W punkcie 2 omówiono związek między tłumieniem pętli abonenckiej na różnych częstotliwościach będący podstawą zaproponowanej metody. W dalszej części opisano opracowany system pomiarowy i rezultaty pomiarów osiągnięte za pomocą tego systemu (punkt 3). Wreszcie punkty 4 i 5 są poświęcone omówieniu zależności między tłumieniem linii a przepływnością pracującego na niej systemu ADSL. Rezultaty pracy podsumowano w punkcie 5.

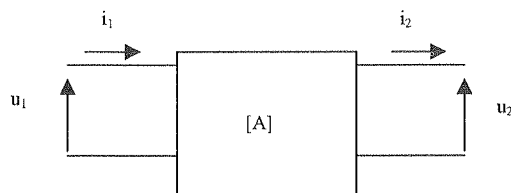
2. ZWIĄZEK MIĘDZY TŁUMIENIEM PĘTLI NA RÓŻNYCH CZĘSTOTLIWOŚCIACH

Dla określenia tłumienności wtrąceniowej danej pętli abonenckiej niezbędne jest wprowadzenie pojęcia macierzy łańcuchowej, A , wiążącej napięcie i prąd wejściowy z napięciem i prądem na wyjściu czwórnika. Pokazano to na rys. 1. Definicja tej macierzy jest następująca

$$\begin{bmatrix} u_1 \\ i_1 \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \cdot \begin{bmatrix} u_2 \\ i_2 \end{bmatrix} = [A] \cdot \begin{bmatrix} u_2 \\ i_2 \end{bmatrix} \quad (1)$$

Nietrudno wykazać, że tłumienność wtrąceniowa Att pomiędzy źródłem (generatorem) o rezystancji R_0 a obciążeniem o takiej samej rezystancji, połączonymi przez czwórnik o macierzy łańcuchowej A (patrz rys. 2), wyraża się zależnością

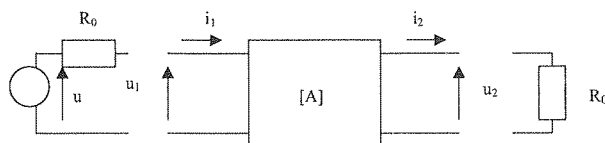
$$Att = 20 \log_{10} \left| a_{11} + a_{22} + \frac{a_{12}}{R_0} + a_{21}R_0 \right| - 6 \quad [\text{dB}] \quad (2)$$



Rys. 1. Napięcia i prądy wykorzystywane w definicji macierzy łańcuchowej

Fig. 1. Voltages and currents used in the chain matrix definition

Przy połączeniu łańcuchowym czwórników macierz łańcuchowa połączonego układu jest równa iloczynowi macierzy łańcuchowych jej elementów składowych, co pozwala w stosunkowo łatwy sposób policzyć tłumienność wtrąceniową takiego połączenia.



Rys. 2. Układ do określenia tłumienności wtrąceniowej

Fig. 2. Setup used for the definition of the insertion loss

Na rys. 3. pokazano pętle zgodne z ETSI, które odpowiadają pętlom abonenckim występującym w Polsce. Z rysunku wynika, że dla policzenia tłumienności wtrąceniowych takich pętli należy określić macierze łańcuchowe: odcinka linii długiej oraz macierz łańcuchową i impedancję wejściową rozwartego na końcu odcinka linii długiej.

Macierz łańcuchowa odcinka linii długiej o długości L jest określona zależnością [8]

$$A = \begin{bmatrix} \cosh(\gamma L) & Z_f \sinh(\gamma L) \\ \frac{\sinh(\gamma L)}{Z_f} & \cosh(\gamma L) \end{bmatrix} \quad (3)$$

Tutaj $\gamma = \alpha + j\beta$ jest współczynnikiem propagacji, zaś Z_f impedancją falową. Przy założeniu zerowej upływności ($G = 0$) oraz pozostałych parametrach jednostkowych linii

równych R , L , C , powyższe parametry dane są następującymi zależnościami ogólnymi [1]

$$\alpha = \frac{1}{\sqrt{2}} \sqrt{\omega C \sqrt{R^2 + \omega^2 L^2} - \omega^2 LC} \quad (4)$$

$$\beta = \omega \sqrt{\frac{LC}{2}} \sqrt{1 + \sqrt{1 + \frac{R^2}{\omega^2 L^2}}} \quad (5)$$

$$Z_f = \frac{\beta}{\omega C} - j \frac{\alpha}{\omega C} \quad (6)$$

Przy tym w badanym zakresie częstotliwości (do ≈ 1.1 MHz) jedynie pojemność jednostkowa linii C nie jest funkcją częstotliwości f , natomiast zarówno R jak i L są od niej zależne [2].

Macierze łańcuchowe czwórników pokazanych na rys. 4 są określone następującymi zależnościami (rys. 4A)

$$B = \begin{bmatrix} 1 & 0 \\ 1/Z_{in} & 1 \end{bmatrix} \quad (7)$$

$$C = \begin{bmatrix} 1 & Z \\ 0 & 1 \end{bmatrix} \quad (8)$$

Z kolei, ważna w dalszych obliczeniach, wartość impedancji wejściowej Z_{in} linii długiej zakończonej impedancją Z_2 wyraża się wzorem [8]

$$Z_{in} = Z_f \frac{Z_2 + Z_f \operatorname{tgh}(\gamma L)}{Z_f + Z_2 \operatorname{tgh}(\gamma L)} \quad (9)$$

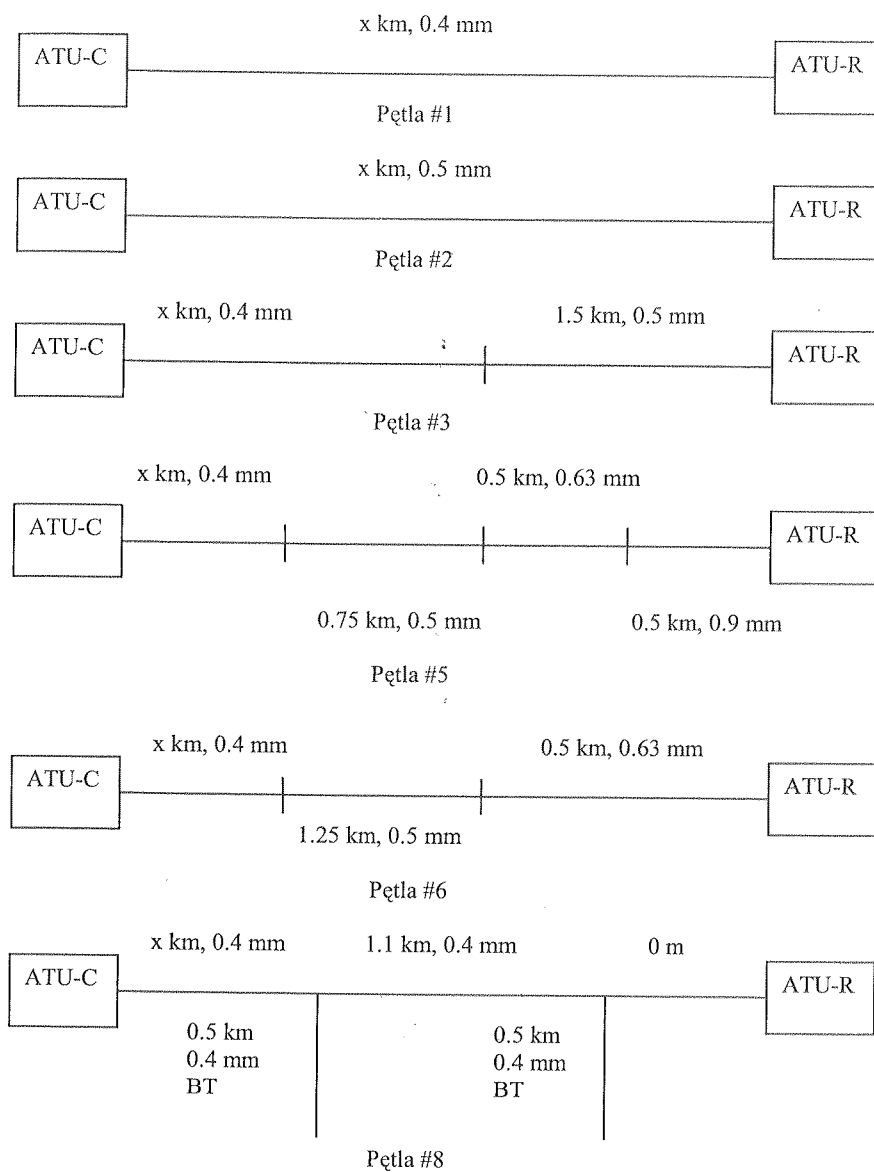
Przy rozwarciu linii na jej końcu ($Z_2 = \infty$) powyższa zależność redukuje się do

$$Z_{in} = Z_f \operatorname{ctgh}(\gamma L) \quad (10)$$

Oznaczając przez $A(y, x)$ macierz łańcuchową odcinka przewodu o średnicy żyły y mm oraz długości x , oraz przez $B(y, x)$ macierz łańcuchową rozwartego odczepu o średnicy żyły y mm i długości x (patrz zależności (7) i (10)) można uzyskać wzory na macierze łańcuchowe poszczególnych pętli abonentkich pokazanych na rys. 3. Podsumowano to w tabl. 1. Należy w tym miejscu zauważyć, że wyznaczniki macierzy A , B i C (3), (7), (8) są równe jedności;

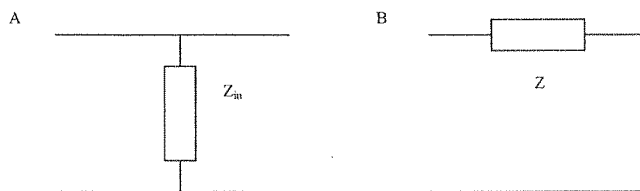
$$\det A = \det B = \det C = 1 \quad (11)$$

Świadczy to o ich odwracalności (symetryczności energetycznej). Również połączenie łańcuchowe takich czwórników jest energetycznie symetryczne (np. $\det A \cdot B = \det A \cdot \det B = 1$). Zatem tłumienność wtrąceniowa (2) dowolnego połączenia łańcuchowego takich czwórników (a więc dowolnej pętli abonentkiej) nie zależy od tego, z którego końca jest mierzona. Można to wykorzystać przy jej pomiarach.



Rys. 3. Pętle ETSI, które odpowiadają pętlom abonenckim występującym w Polsce

Fig. 3. ETSI loops used in Poland: ATU-C and ATU-R terminal units at the central office and remote site, respectively, BT-bridged tap



Rys. 4. Czwórnik stosowane w liniach transmisyjnych (np. przy pupinizacji)

Fig. 4. Four pole devices used in the transmission lines

Tabela 1

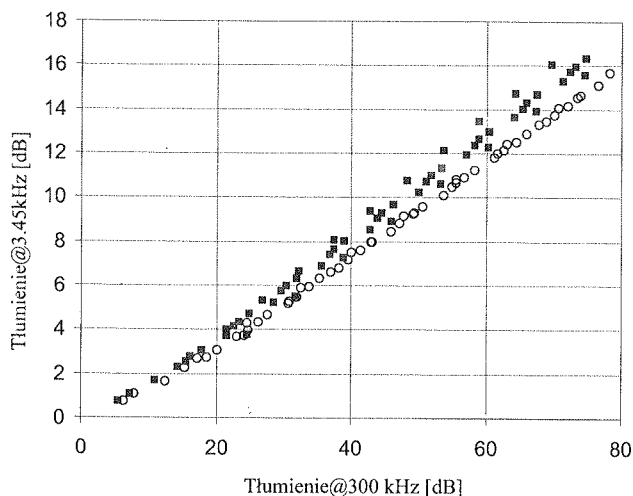
Macierze łańcuchowe poszczególnych pętli ETSI z rys. 3
Chain matrices of the ETSI loops used in Poland

Typ pętli	Macierz łańcuchowa
#1	$A(0.4, x)$
#2	$A(0.5, x)$
#3	$A(0.4, x). A(0.5, 1.5)$
#5	$A(0.4, x). A(0.5, 0.75). A(0.63, 0.5). A(0.9, 0.5)$
#6	$A(0.4, x). A(0.5, 1, 25). A(0.63, 0.5)$
#8	$A(0.4, x). B(0.4, 0.5). A(0.4, 1.1). B(0.4, 0.5)$

Po obliczeniu wypadkowej macierzy łańcuchowej pętli abonenckiej można zgodnie z zależnością (2) obliczyć tłumienność wtrąceniową tej pętli. Obliczenia te przeprowadzono przy zmieniających się co 500 m odległościach x dla różnych częstotliwości z zakresu pasma telefonicznego ($R = 600\Omega$) oraz na częstotliwości 300 kHz, na której zazwyczaj określa się tłumienie linii w systemach ADSL ($R = 100\Omega$). Wyniki obliczeń wykazują, że związek między obydwoma tłumiennościami jest tym wyraźniejszy im większa jest częstotliwość z zakresu pasma telefonicznego. Dalej zaprezentowano wyniki uzyskane dla częstotliwości 3.45 kHz, bliskiej górnemu krańcowi pasma telefonicznego, a jednocześnie będącą jedną z częstotliwości sondowania linii w protokole V.34 [3]. Natomiast częstotliwość ta może być nieco za wysoka z praktycznego punktu widzenia, gdyż na tej częstotliwości na tłumienie linii mogą mieć wpływ transmitancje filtrów nadawczego i odbiorczego [13].

Obliczenia wykonano dla dwóch typów kabli: z izolacją PE oraz z izolacją papierową, zaś dane odnośnie wartości parametrów $RGLC$ (ich zależność częstotliwościowa) wzięto z zalecenia [2].

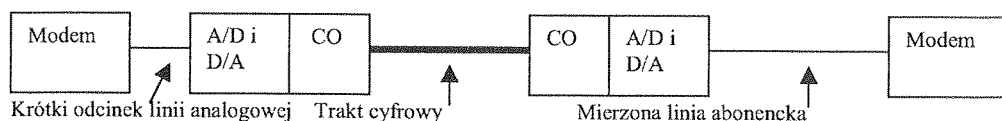
Wyniki obliczeń przedstawia rys. 5, gdzie pokazano tłumienności wtrąceniowe wielu różnych pętli abonenckich na dwóch wspomnianych wyżej częstotliwościach.



Rys. 5. Zależność tłumienności wtrąceniowej różnych pętli abonentkich na częstotliwościach 3.45 kHz i 300 kHz: prostokąty — kable z izolacją PE, kółka — kable z izolacją papierową

Fig. 5. Insertion losses of different ETSI loops at 3.45 kHz and 300 kHz:

o — paper insulation, ■ — PE insulation



Rys. 6. Schemat układu pomiaru parametrów linii abonentckiej przy pomocy pary modemów analogowych. Oznaczenia: A/D i D/A — przetworniki analogowo-cyfrowe i cyfrowo analogowe, CO — centrala telefoniczna

Fig. 6. Set up for subscriber loop measurements by means of pair of analog V.34 modems:

A/D & D/A — respective converters, CO — central office

Z rysunku tego widać, że z dobrym przybliżeniem można przyjąć, iż obydwie tłumienności są ze sobą związane zależnością liniową, którą można przybliżyć przez

$$Att(300 \text{ kHz}) \approx 5 + 4.5 \cdot Att(3.45 \text{ kHz})[\text{dB}] \quad (12)$$

Błąd tego przybliżenia rośnie wraz z tłumieniem linii, ale nawet dla największych jego wartości nie przekracza 5 dB.

Ponieważ tłumienność jednostkowa linii decyduje w zasadniczym stopniu o tłumienności wtrąceniowej, zatem wzór (12) można wyjaśnić rozpatrując zależność tłumienności jednostkowej linii od częstotliwości dla popularnych typów kabli.

Dla niskich częstotliwości f tłumienność jednostkowa (określona w Np/km) wyraża się wzorem [1]

$$\alpha_n \approx \sqrt{\pi f R_n C}, \quad (13)$$

podczas gdy dla wysokich częstotliwości [1]

$$\alpha_w \approx \frac{R_w}{2 \sqrt{L_w/C}} = \frac{R_w \sqrt{C}}{2 \sqrt{L_w}} \quad (14)$$

We wzorach (13), (14) przez R_n i R_w oznaczono rezystancję jednostkową linii, odpowiednio na niskiej i wysokiej częstotliwości, L_w jest indukcyjnością jednostkową linii na wysokiej częstotliwości, zaś f jest częstotliwością (niską).

Ostatecznie mamy

$$\frac{\alpha_w}{\alpha_n} = \frac{R_w}{2 \sqrt{\pi f L_w R_n}} \quad (15)$$

Wartości stosunku $\alpha(300\text{kHz})/\alpha(3.45\text{kHz})$ obliczone zgodnie z wzorem (15) dla różnych typów kabli (dane kabli wzięto z [2]) podaje tabl. 2. Z tabl. 2. widać, że stosunek

Tabela 2

Wielkość stosunku tłumienności jednostkowych $\alpha(300\text{kHz})/\alpha(3.45\text{kHz})$ dla różnych typów kabli

Attenuation ratio $\alpha(300\text{kHz})/\alpha(3.45\text{kHz})$ for various cable types

Izolacja kabla	PE				papier			
Średnica żyły w mm	0.4	0.5	0.63	0.9	0.4	0.5	0.65	0.9
$\alpha(300\text{ kHz})/\alpha(3.45\text{ kHz})$	4.27	3.98	4.18	4.56	4.63	4.63	4.44	4.46

tłumienności na obydwu częstotliwościach mało zależy od średnicy żyły i ma wartość zbliżoną do współczynnika we wzorze (12). Należy wszakże pamiętać, iż zależność (12) opisuje tłumienność wtrąceniową, a nie jednostkową. Niemniej jednak porównując tabl. 2 i rys. 5 widać zgodnie, że dla kabli z izolacją PE wartość stosunku tłumienności (jednostkowych i wtrąceniowych) na częstotliwościach 300 kHz i 3,45 kHz jest nieznacznie mniejsza, a jej wartości mają większy rozrzut, aniżeli dzieje się to dla kabli z izolacją papierową.

Niestety, jeśli do pomiarów parametrów linii w paśmie głosowym wykorzystujemy parę modemów analogowych, tak jak jest to realizowane w zaproponowanym systemie pomiarowym pokazanym na rys. 6, bezpośrednie skorzystanie z zależności (12) lub podobnej nie jest możliwe. Składa się na to kilka przyczyn:

1. Centrale telefoniczne w trakcie cyfrowym pomiędzy modemami pomiarowymi mogą wprowadzać pewne stałe tłumienie sygnału (tzw. digital pad — typowo w zakresie 0 ... 7 dB), które dodaje się do całkowitego tłumienia linii.
2. Pomiar wartości absolutnej poziomu mocy odbieranej przez modem jest obarczony dosyć dużym błędem; znacznie dokładniejszy jest pomiar względny (różnic mocy).
3. Nie wszystkie modemy udostępniają wartości poziomu mocy nadawanych i odbieranych, co jest niezbędne do wyliczenia tłumienia linii.

Z tego względu do praktycznego wykorzystania nadaje się tylko charakterystyka częstotliwościowa kanału (charakterystyka przenoszenia linii), która jest zdejmowana w fazie 2 inicjalizacji połączenia modemowego V.34 [3] (sygnały L1 i L2). W tej fazie jeden z modemów generuje sygnały harmoniczne o częstotliwościach 150, 300, 450...3600, 3750 Hz, co 150 Hz, z wyjątkiem (być może) niektórych tonów, a drugi mierzy poziom tych sygnałów. W ten sposób zdejmowana jest charakterystyka przenoszenia linii analogowej, przy czym, jeśli linia analogowa w centrali nie wnosi zniekształceń (linia ta jest krótka — rys. 6), to pomiar dotyczy wyłącznie pętli abonenckiej. Wszelkie stałe tłumienie jakiego doznaje sygnał w torze cyfrowym nie ma wpływu na wyniki pomiaru (w tym sensie, że zarówno pasmo, jak i nachylenie krzywej tłumienia nie zależy od tego stałego tłumienia). Z wymienionych powyżej przyczyn (digital pad(s), niewiarygodne odczyty mocy absolutnej z modemu) do wykorzystania nadają jedynie względne zmiany mocy i to jedynie w (maksymalnie dużym) paśmie, na który nie mają wpływu charakterystyki filtrów nadawczych i odbiorczych. Dobrym wyborem takiego pasma jest zakres 600...3000 Hz. W tym zakresie z danych udostępnianych przez modem obliczane jest nachylenie charakterystyki tłumienia, przy czym wykorzystywana jest średniokwadratowa aproksymacja tej charakterystyki linią prostą. Nachylenie charakterystyki tłumienia jest definiowane jako wyrażone w dB nachylenie tej prostej odniesione do przedziału 600-3000 Hz. Takie zdefiniowanie parametru, na podstawie którego obliczane jest tłumienie linii, ogranicza błędy związane z kwantowaniem wyników pomiarowych oraz ewentualne błędy związane z wpływem filtrów nadawczego i/lub odbiorczego. Jeśli oznaczymy przez x_i wyrażoną w kHz częstotliwość pomiarową z zakresu (0.6...3) kHz, zaś przez y_i — odpowiadający jej wynik pomiaru poziomu sygnału, to szukane nachylenie prostej N_A , zgodne z aproksymacją średniokwadratową dane jest zależnością

$$N_A = \frac{XY - X \cdot Y}{X_2 - (X)^2} \cdot 2.4 \quad [\text{dB}] \quad (16)$$

gdzie

$$X = \frac{1}{N} \sum_{i=1}^N x_i \quad (17)$$

$$Y = \frac{1}{N} \sum_{i=1}^N y_i \quad (18)$$

$$XY = \frac{1}{N} \sum_{i=1}^N x_i y_i \quad (19)$$

$$X_2 = \frac{1}{N} \sum_{i=1}^N x_i^2 \quad (20)$$

zaś N jest liczbą punktów pomiarowych (częstotliwości). W przypadku bezbłędnych danych $N = 17$.

Aby odnieść obliczone nachylenie N_A do szukanej wartości tłumienia linii na 300 kHz $Att(300\text{kHz})$ dokonano obliczeń par wartości N_A , $Att(300\text{kHz})$ dla różnych typów pętli abonenckich, przy różnych typach izolacji, średnicach żyły i długościach. Przy znanym typie kabla/pętli wartość $Att(300\text{kHz})$ jest obliczana na podstawie N_A poprzez interpolację liniową danych z tablicy zawierającej pary N_A , $Att(300\text{kHz})$ odpowiadające danemu kablowi/pętli. Jeśli dane dotyczące pętli są niepełne (np. znany jest tylko typ izolacji, lub jedynie średnica żyły) lub ich brak, to dokonuje się interpolacji we wszystkich tablicach odpowiadających tym niepełnym danym. Dla przykładu przy znanej izolacji, np. PE, interpolację przeprowadza się we wszystkich pętlach używających kabla o takiej izolacji. Wynikowe tłumienie $Att(300\text{kHz})$ jest obliczane jako średnia arytmetyczna wyników interpolacji dla poszczególnych tabel.

3. POMIARY TŁUMIENNOŚCI

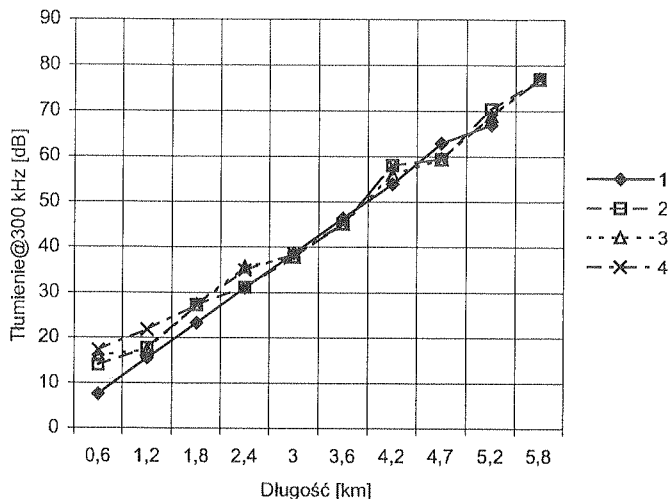
Analiza związku między tłumiennością wtrąceniową pętli abonenckich dla różnych częstotliwości nie byłaby możliwa bez weryfikacji praktycznej (pomiarów). W tym celu stworzony został odpowiedni system pomiarowy. Składa się on z dwóch części: **centralowej** — serwera pomiarowego umieszczonego w centrali telefonicznej jak najbliżej linii telefonicznych,

abonenckiej — modułów wyniesionych (komputery osobiste) znajdujących się na końcach mierzonych linii telefonicznych.

Serwer pomiarowy i moduły wyniesione posiadają wbudowane modemy telefoniczne, które służą do podłączenia się na obu końcach mierzonej linii telefonicznej. Podział systemu na część centralową i abonencką oraz stworzone odpowiednie oprogramowanie kontrolno-pomiarowe umożliwiają automatyczne szacowanie tłumienności linii telefonicznej (pętli abonenckiej) na częstotliwości 300 kHz.

Dla sprawdzenia przydatności zaproponowanej metody dokonano w dn. 14.04.2004 w CBR TP S.A. pomiarów tłumienności kabla (średnica żyły 0.4 mm, izolacja PE, żelowany; typ XzTKMXw 15x4x0.4 firmy Telefonika) o zmiennej długości (długość zmienna od 0 do 5.8 km, w krokach co 0.6 km). Para modemów była dołączona krótkimi kablami z lokalną centralą telefoniczną, przy czym w jedno z połączeń wtrącony był wspomniany kabel o zmiennej długości. Zależność tłumienności tego kabla od długości i częstotliwości zmierzona została również niezależnie innym przyrządem (analizator widma). Istotne jest to, że przy zerowej długości mierzonego kabla wystąpiło pewne niezerowe nachylenie charakterystyki częstotliwościowej mierzonej parą modemów (w opisywanym przypadku $N_A = 0.82$). Było one spowodowane łączną charakterystyką toru (filtry nadawcze i odbiorcze w modemach, centrali, niezerowa długość kabli doprowadzających itp.). Ponieważ wartość nachylenia była stosunkowo duża zachodziła konieczność korekcji uzyskanych rezultatów o tę wartość (odjęcie $N_A = 0.82$ od wartości nachyleń obliczonych z danych z modemu). Skorygowane wyniki porównane są z rzeczywistą charakterystyką tłumienności. Rezultaty pokazano na rys. 7. Widać dobrą zgodność obydwu grup danych z wyjątkiem zakresu małych tłumienności, gdzie po-

miar nachylenia jest mało dokładny, a małe nachylenia odpowiadają stosunkowo dużym tłumieniom. W omawianym przypadku korekcja była prosta, jednak dla nieznanych a priori linii i central należy dokonać pewnej korekcji statystycznej, np. mierząc charakterystyki pętli o zerowych długościach z kilku różnych central i uśredniając wyniki.

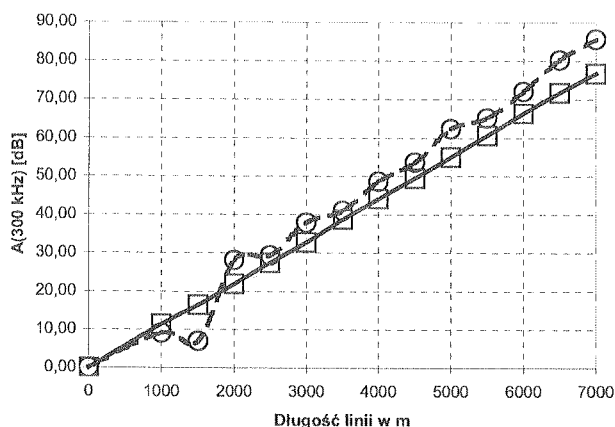


Rys. 7. Tłumienność @ 300 kHz pary w kablu (0.4 mm, PE) w funkcji jego długości w km (pomiar w CBR Warszawa): 1 — pomiar rzeczywisty, 2,3,4 — wyliczone z pomiarów modemowych dla 3 różnych par

Fig. 7. The loop attenuation (twisted pair, 0.4 mm, PE insulation) as a function of its length (measurement in Warsaw): 1— real value, 2, 3, 4 — calculated from the modems results for three different pairs

Dalszej weryfikacji pomiarowej proponowanej metody dokonano w listopadzie 2004 r. w laboratorium pomiarowym Telekomunikacyjnej Izby Pomiarowej w Krakowie [14]. Urządzenie pomiarowe znajdowało się w CBR w Warszawie natomiast badana linia oraz modem abonencki we wspomnianym już laboratorium pomiarowym w Krakowie. Modem abonencki poprzez rzeczywistą linię poligonu kablowego (miedziana para symetryczna) o zmiennej długości z zakresu 0...7 km był dołączony do lokalnej centrali w Krakowie. Modem pomiarowy dołączony był do centrali w Warszawie przez krótką linię miedzianą o pomijalnej tłumienności wtęceniowej. Podczas pomiarów modem abonencki (Kraków) łączył się wielokrotnie z urządzeniem pomiarowym w Warszawie. Wcześniej dokonano pomiaru rzeczywistych tłumienności wszystkich badanych pętli abonenckich na częstotliwości 300 kHz za pomocą zestawu: generator i miernik poziomu. Ponadto wykonano pomiary przy dołączonych do modemu abonenckiego dwóch aparatów telefonicznych, jak również pomiary w obecności szumu przenikowego. Jednakże w tym ostatnim przypadku dołączenie źródła szumów prowadziło do zmian impedancji w układzie modemu i dlatego te ostatnie rezultaty nie będą tu prezen-

owane. Wyniki pozostałych pomiarów pokazano na rys. 8-10. Uzyskano tu podobną zgodność jak w pomiarach warszawskich (rys. 7), przy czym ponownie największym problemem było niezerowe nachylenie charakterystyki przenoszenia w paśmie głosowym przy zerowej długości pętli abonenckiej spowodowane filtrami w modemach i centrali. Podobnie jak przy pomiarach warszawskich konieczna była korekcja wyników o to nachylenie N_A . Zaprezentowane na rys. 8-10 rezultaty pomiarów uwzględniają dokonaną korekcję. Na uwagę zasługuje bardzo niewielki wpływ na pomiary dołączenia aparatów telefonicznych od strony abonenckiej (rys. 10).



Rys. 8. Wyniki pomiarów tłumienności wtrąceniowej A (300 kHz) [dB] (oś pionowa) w funkcji długości linii L [m] (oś pozioma) przeprowadzonych w Laboratorium Pomiarowym Telekomunikacyjnej Izby Pomiarowej w Krakowie: linia #1 bez zakłóceń, linia ciągła — wartość dokładna, linia przerywana — wyniki pomiaru modemem (uśrednione za 6 do 8 pomiarów)

Fig. 8. The loop attenuation at 300 kHz in dB as a function of its length in m (measurements in Kraków), LOP#1: solid line — real value, dashed line — modem results (averaged over 6 to 8 measurements)

4. ZWIĄZEK MIĘDZY TŁUMIENIEM LINII A JEJ PRZEPUSTOWOŚCIĄ

4.1. ESTYMACJA TŁUMIENIA LINII W FUNKCJI CZĘSTOTLIWOŚCI

Dla określenia przepustowości systemów ADSL potrzebne jest określenie zależności częstotliwościowej tłumienia w całym paśmie ADSL (26...1104 kHz) przy zadanej wartości tłumienności pętli na częstotliwości 300 kHz: $Att(300 \text{ kHz})$. W tym celu można skorzystać z ogólnej zależności funkcyjnej tłumienia pary symetrycznej od częstotliwości (pulsacji) [1], danej w postaci

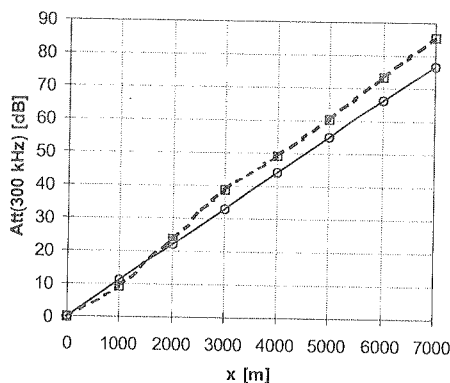
$$\alpha = A + B\sqrt{\omega} + C\omega \quad (21)$$

Rys.
przepr
linia #2

Fig.

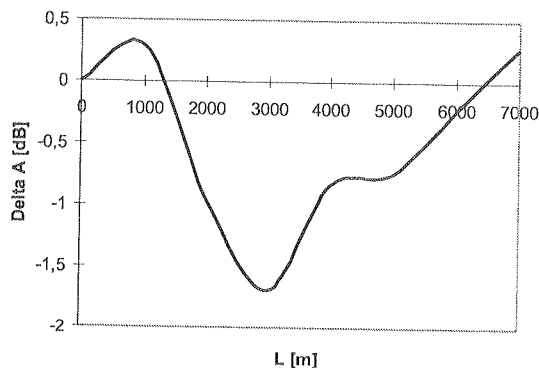
Rys.
dołączeni
wskaza

Fig. 10. A



Rys. 9. Wyniki pomiarów tłumienności wtrąceniowej A (300 kHz) w funkcji długości linii L przeprowadzonych w Laboratorium Pomiarowym Telekomunikacyjnej Izby Pomiarowej w Krakowie: linia #2 bez zakłóceń, linia ciągła — wartość dokładna, linia przerywana — wyniki pomiaru modemem (uśrednione za 3 pomiary)

Fig. 9. The loop attenuation at 300 kHz in dB as a function of its length in m (measurements in Kraków), Loop#2: solid line — real value, dashed line — modem results (averaged over 3 measurements)



Rys. 10. Różnica tłumienności Delta A (@300 kHz) w pomiarach modemowych spowodowana dołączeniem do końca abonentkiego dwóch aparatów telefonicznych. Różnica między obliczonymi ze wskazań modemów tłumiennościami linii z aparatami telefonicznymi i bez tych aparatów w funkcji długości linii (linia #2)

Fig. 10. Attenuation difference in modem measurements caused by connection of two telephone sets to the line at the subscriber premises as a function of the line length

gdzie A, B, C są pewnymi stałymi. W naszym przypadku zależność ta przyjmuje postać

$$Att(f) = Att(300kHz) \left[a + b \sqrt{\frac{f}{300kHz}} + c \frac{f}{300kHz} \right] \quad [dB] \quad (22)$$

gdzie stałe a, b, c powinny być dobrane dla uzyskania jak najlepszej zgodności z danymi doświadczalnymi.

Porównanie zależności tłumienia linii od częstotliwości obliczonej z wzoru (22), w którym przyjęto $a = 0.5, b = 0.3, c = 0.2$, z rzeczywistymi tłumieniami różnych pętli abonenckich uzyskanymi z [2] pokazano na rys. 11, 12. Widać, że zgodność między wartością teoretyczną (22) a wartościami rzeczywistymi jest bardzo dobra, o ile pętla nie zawiera odczepów. W przypadku odczepów (pętla C#4, C#7, T#9, T#13, [2]) największe rozbieżności występują wtedy, gdy częstotliwość 300 kHz znajduje się w pobliżu jednego z rezonansów pętli. Widać to dobrze zwłaszcza w przypadku pętli T#9 i T#13, gdzie rzeczywista pętla wykazuje wzrost tłumienia, co powoduje przesunięcie całej krzywej tłumienia w górę. Należy jednak pamiętać, że w naszym przypadku tłumienie na częstotliwości 300 kHz jest uzyskiwane przez ekstrapolację wyników pomiaru na niższych częstotliwościach, co pozwala się uniezależnić od ewentualnych rezonansów w pętlach z odczepami.

Rys. 11

Fig.

4.2. ESTYMACJA PRZEPŁYWNOŚCI LINII W OBECNOŚCI ZAKŁÓCEŃ

Tłumienie linii łącznie z typem występujących zakłóceń określa stosunek sygnału do szumu w odbiorniku, a zatem również i przepływność binarną możliwą do osiągnięcia w systemie ADSL pracującym w danej pętli abonenckiej. W tym miejscu należy zwrócić uwagę na rozróżnienie różnych typów przepływności:

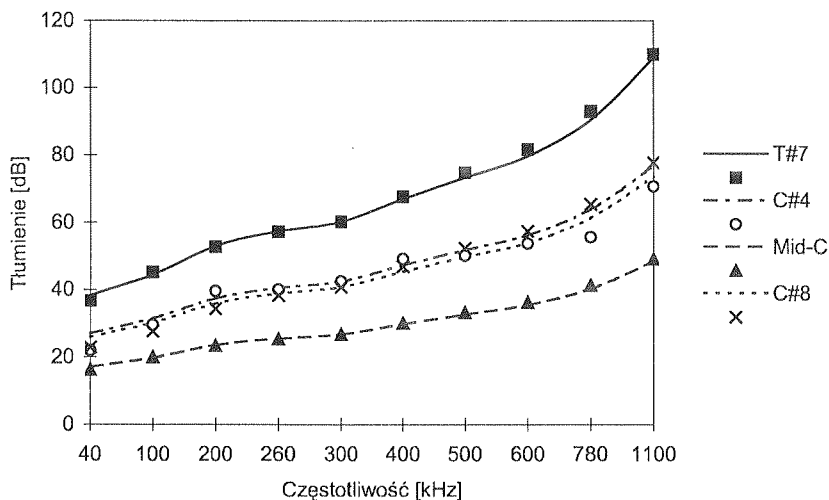
- przepływność linii (ang. line rate) uwzględnia również dane pochodzące z nagłówków, kodowania nadmiarowego (Reeda-Solomona, kratowego) itd.
- rzeczywista przepływność danych użytkownika (ang. net data rate, payload bit rate) jest mniejsza i może wynosić jedynie około 75% przepływności linii.

Zależności przepływności binarnej od tłumienia linii na częstotliwości 300 kHz przy różnych typach zakłóceń zebrano na rys. 13–15. Pokazane wartości dotyczą przepływności linii (dane dotyczące przepływności danych użytkownika odniesiono do przepływności linii). Dane opracowano na podstawie minimalnych wartości przepływności przy danych typach pętli i zakłócenia zawartych w odpowiednich zaleceniach [2,4,6] oraz pomiarów rzeczywistych modemów pracujących w różnych warunkach [5,7,15]. Wyróżniono przy tym zakłócenia typu T1/E1 oraz ETSI-B ze względu na ich wpływ silnie ograniczający przepływność systemu ADSL. Z rysunków widać, że dla tych samych tłumień linii przepływności dla zakłóceń T1/E1 oraz ETSI-B są generalnie mniejsze niż dla pozostałych typów zakłóceń.

Rys. 12.

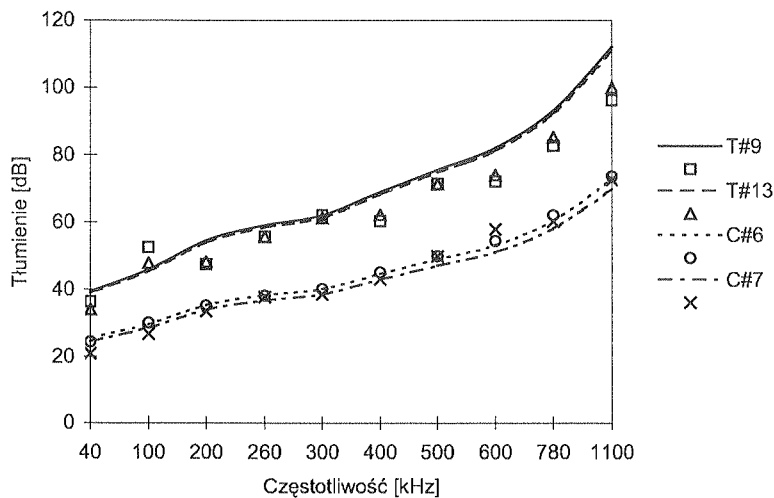
Inna możliwość oszacowania spodziewanej przepływności przy danym tłumieniu linii polega na analizie większej grupy danych pomiarowych określających uzyskana przepływność przy danym tłumieniu linii na częstotliwości 300 kHz. Takie dane są

Fig. 1



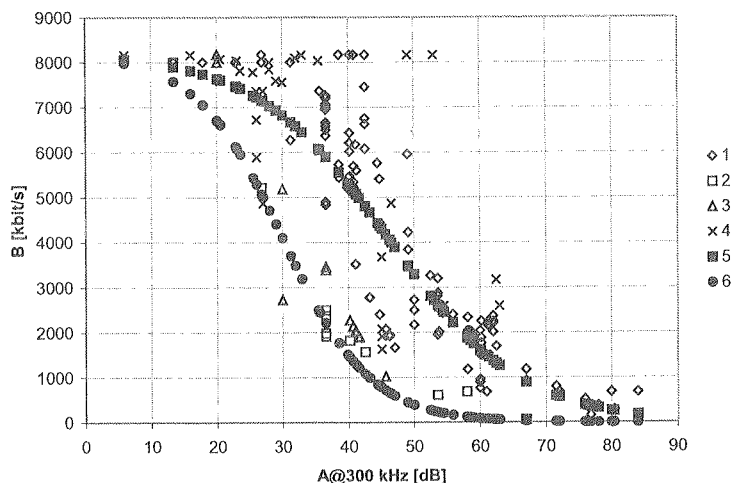
Rys. 11. Zależność tłumienia linii określonego przez model (22) — linie oraz odpowiadającego mu tłumienia rzeczywistych pętli abonenckich — punkty, dla typów pętli określonych po prawej stronie wykresu

Fig. 11. Loop attenuation versus frequency for the loop types listed at the right side: lines — attenuations according to the model, points — actual loop attenuations



Rys. 12. Zależność tłumienia linii określonego przez model (22) — linie oraz odpowiadającego mu tłumienia rzeczywistych pętli abonenckich — punkty, dla typów pętli określonych po prawej stronie wykresu

Fig. 12. Loop attenuation versus frequency for the loop types listed at the right side: lines — attenuations according to the model, points — actual loop attenuations



Rys. 13. Zależność przepływności linii w kierunku do abonenta B (downstream) od tłumienia linii na częstotliwości 300 kHz dla różnych typów zakłóceń: 1 — wszystkie typy zakłóceń z wyjątkiem zakłóceń T1/E1 i ETSI-B, 2 — zakłócenia T1/E1, 3 — zakłócenia ETSI-B, 4 — pomiary rzeczywistych systemów w warunkach eksploatacji [7], 5 — przybliżenie funkcją sigmoidalną (23) o najmniejszym błędzie średniokwadratowym ($p = 46, k = 0.1$), 6 — przybliżenie funkcją sigmoidalną (23) ($p = 30, k = 0.15$) zapewniająca w przybliżeniu 98% prawdopodobieństwo, że rzeczywista przepływność nie będzie mniejsza od wyznaczonej z wzoru (23)

Fig. 13. Downstream line bit rate, B, against the line attenuation at 300 kHz for various interference types. Recommendations: 1 — all interferences except T1/E1 and ETSI-B, 2 — T1/E1 interference, 3 — ETSI-B interference. 4 — Measurement of real ADSL systems. 5 — Approximation with the function (23) with the minimum mean square error ($p = 46, k = 0.1$). 6 — Approximation with the function (23) that ensures with 98% probability that the actual bit rate will not be less than calculated ($p = 30, k = 0.15$)

pokazane na rys. 13–15. Wynika z nich, że dobrym przybliżeniem zależności spodziewanej przepływności od tłumienia linii na 300 kHz jest funkcja sigmoidalna znana z teorii sieci neuronowych i zbiorów rozmytych. Warto tutaj dodać, że przybliżenie zależności przepływności od tłumienia za pomocą innych zależności funkcyjnych (np. przeskalowany arc ctg) prowadziło do większych błędów średniokwadratowych. Funkcja sigmoidalna w omawianym przypadku ma postać

$$y = \frac{B_{MAX}}{1 + \exp[k(x - p)]} \quad [\text{Mbit/s}] \quad (23)$$

gdzie B_{MAX} jest maksymalną przepływnością w danym kierunku transmisji ($B_{MAX} \approx 8 \text{ Mbit/s down, } B_{MAX} \approx 880 \text{ kbit/s up}$), y jest spodziewana przepływność, x tłumieniem linii na 300 kHz, zaś k, p są parametrami, które należy wybrać w zgodzie z danymi doświadczalnymi. Możliwe są tutaj dwa podejścia. Pierwsze z nich polega na takim wy-

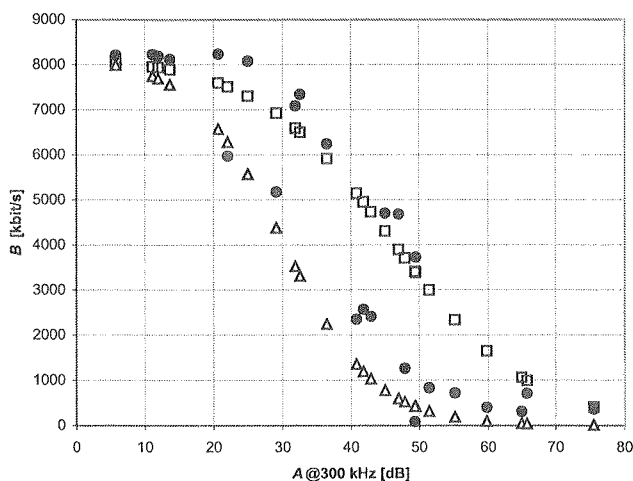
borze p
pomiar
downst
przyjm
z rys.
odchyl
określ
zaznac
giczne
wartoś
 $p = 52$
zauwa
silnie
pokaza

Rys.
rzeczy
średnic

Fig.
op
functi

W wi
racze

borze parametrów k , p , aby zminimalizować błąd średniokwadratowy między danymi pomiarowymi a funkcją (23). Zależność tego błędu od parametrów k , p dla transmisji downstream pokazano na rys. 16. Minimalna wartość odchylenia standardowego jest przyjmowana dla $k \approx 0.1$, $p \approx 46$ (wtedy odchylenie wynosi 1.57 Mbit/s dla danych z rys. 13) przyjęcie takiej samej aproksymacji dla danych kontrolnych z rys. 14, dało odchylenie wynoszące 1.47 Mbit/s (dane kontrolne nie były brane pod uwagę przy określaniu parametrów krzywej sigmoidalnej). Na obydwu wspomnianych rysunkach zaznaczono również aproksymacje odpowiednimi funkcjami sigmoidalnymi. Analogiczne dane dla kierunku transmisji upstream pokazano na rys. 15. W tym przypadku wartości parametrów k , p dające minimalny błąd średniokwadratowy są równe $k = 0.1$, $p = 52$, zaś odpowiadające im odchylenie standardowe wynosi 93 kbit/s. Warto tutaj zauważyć, że odchylenie standardowe danych pomiarowych od funkcji sigmoidalnej silnie zależy od tłumienia linii i jest największe dla jego środkowych wartości, co pokazano na rys. 17.

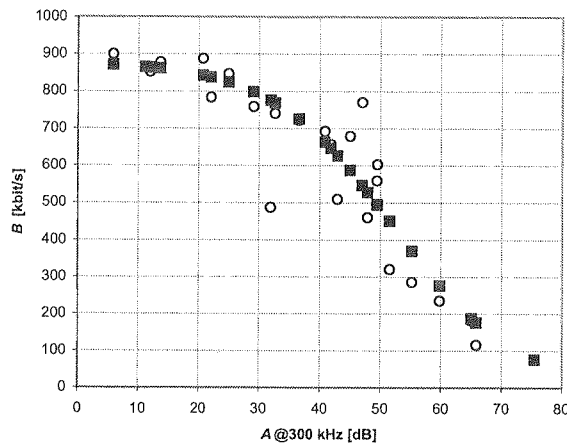


Rys. 14. Zależność przepływności downstream B od tłumienia linii A na 300 kHz dla pomiarów rzeczywistych systemów ADSL pracujących na kablach o średnicy żyły 0.5 mm [15] (dane kontrolne — kółka) oraz aproksymacja tej zależności funkcjami sigmoidalnymi (23): minimalny błąd średniokwadratowy ($k = 0.1$, $p = 46$) — kwadraty, oraz 98% prawdopodobieństwo nie popełnienia błędu ($k = 0.15$, $p = 30$) — trójkąty

Fig. 14. Downstream line bit rate, B , against the line attenuation at 300 kHz for real ADSL systems operating over various 0.5 mm twisted pair cables [15] (circles), and the approximation with the function (21): the minimum mean square error ($p = 46$, $k = 0.1$) — (squares); 98% probability that the actual bit rate will not be less than calculated ($p = 30$, $k = 0.15$) (triangles)

W większości przypadków operatorów sieci interesuje jednak nie zależność średnia, a raczej minimalna wartość przepływności, która z bardzo dużym prawdopodobieństwem

(np. 98%) jest co najmniej osiągnięta dla danej wartości tłumienia. Takie zależności mogą być również wyrażone funkcjami sigmoidalnymi, jednak o innych parametrach (w omawianym przypadku $k = 0.15$, $p = 30$) od tych, które minimalizują błąd średniokwadratowy. Funkcje te pokazano na rys. 13 i 14. Patrząc na dane kontrolne widać (rys. 14), że uzyskane funkcje sigmoidalne są z nimi zgodne.



Rys. 15. Zależność przepływności linii B w kierunku do centrali (upstream) od tłumienia linii na częstotliwości 300 kHz dla pomiarów rzeczywistych systemów ADSL pracujących na kablach o średnicy żyły 0.5 mm [15] — kółka oraz średniokwadratowe przybliżenie tej zależności funkcją sigmoidalną (23) zapewniającą minimalny błąd średniokwadratowy ($k = 0.1$, $p = 52$, odchylenie standardowe = 93 kbit/s) — kwadraty

Fig. 15. Upstream line bit rate, B , against the line attenuation at 300 kHz for real ADSL systems operating over various 0.5 mm twisted pair cables [15] (circles), and the approximation with the function (23) with the minimum mean square error ($p = 52$, $k = 0.1$, standard deviation = 93 kbit/s) — squares

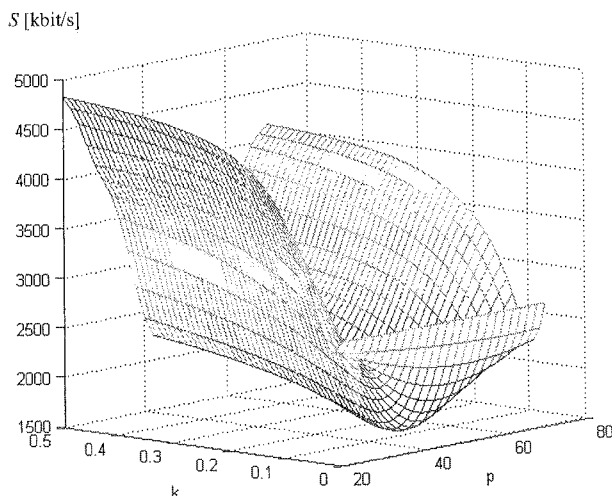
4.3. TEORETYCZNA ZALEŻNOŚĆ PRZEPŁYWNOŚCI LINII OD JEJ TŁUMIENIA

Spróbujemy teraz teoretycznie powiązać przepływność w systemie ADSL z tłumieniem linii na częstotliwości 300 kHz.

Przy założeniu modulacji DMT(QAM), rozsądnych wartości zysku kodowania i marginesu szumowego, przyjęciu elementowej stopy błędów $BER = 10^{-7}$, oraz odpowiedniego algorytmu alokacji mocy w poszczególnych kanałach, wartość przepływności B (przepływność linii) w systemie ADSL może być określona przez [10-12]

$$B \approx m \cdot \sum_i \log_2 \left(1 + \frac{SNR_i}{10} \right) \quad [bit/s] \quad (24)$$

Tutaj m jest szybkością transmisji w symbolach/s ($m = 4000$ 1/s [9] — rzeczywista szybkość transmisji wynosi $69/68 \cdot 4000$ 1/s, gdyż co 68 ramek z danymi wstawia się jedną ramkę synchronizacyjną), sumowanie rozciąga się na wszystkie aktywne kanały, zaś SNR_i jest wartością stosunku sygnału do szumu w kanale o indeksie i . Wzór (24) nie uwzględnia tego, że w rzeczywistym systemie ADSL maksymalna liczba symboli w każdym kanale jest ograniczona do $2^8 \dots 2^{15}$ [9]. Może to prowadzić do błędów oszacowania wyższych przepływności sumarycznych. Zamieniając w (24) logarytmy



Rys. 16. Zależność średniego odchylenia standardowego S [kbit/s] rzeczywistych przepływności downstream od zależności przepływności od tłumienia @ 300 kHz wyrażonej funkcją sigmoidalną (23) przy różnych parametrach k, p tej funkcji. Minimalna wartość odchylenia standardowego jest przyjmowana dla $k=0.1, p=46$ (wtedy $S = 1.57$ Mbit/s dla danych ogólnych; przyjęcie takiej samej aproksymacji dla danych kontrolnych dało wynik $S = 1.47$ Mbit/s)

Fig. 16. Mean standard deviation of the actual downstream bit rates from the sigmoidal function (21) for various k, p parameter values

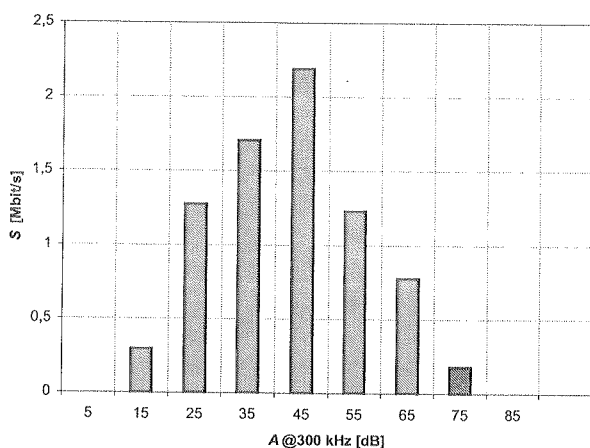
dwójkowe na dziesiętne oraz przyjmując, że w każdym kanale $SNR_i \gg 10$ otrzymujemy

$$B \approx \frac{m}{3} \sum_i 10 \log_{10} (SNR_i/10) = \frac{m}{3} \left(\sum_i SNR_{i\text{dB}} - i \cdot 10_{\text{dB}} \right) = \frac{m}{3} \left(\sum_i P_i - \sum_i N_i - i \cdot 10_{\text{dB}} - \sum_i Att(f_i) \right) \quad [\text{bit/s}] \quad (25)$$

W powyższym wzorze P_i jest mocą nadawaną w i -tym kanale, w dBm, N_i jest mocą zakłóceń w odbiorniku w i -tym kanale, w dBm, zaś $Att(f_i)$ jest tłumieniem linii na częstotliwości i -tego kanału. Skorzystano tutaj z oczywistej zależności

$$SNR_{i\text{dB}} = P_i - Att(f_i) - N_i \quad (26)$$

Jeśli przybliżymy zależność częstotliwościową tłumienia linii jako (22)



Rys. 17. Odchylenie standardowe S rzeczywistych wartości pomiarów przepływności od krzywej sigmoidalnej minimalizującej błąd średniokwadratowy w zależności od wartości tłumienia linii (na wykresie podane są środki przedziałów, szerokość przedziału wynosi 10 dB, przedziały w których znajdowało się mniej jak 5 wyników pomiarów nie zostały uwzględnione)

Fig. 17. Standard deviation of the actual downstream bit rates from the sigmoidal function minimizing the mean square error as a function of line attenuation (central attenuations are shown, the intervals with less than 5 results were discarded)

$$Att(f_i) = Att(300kHz) \cdot \left(a + b \sqrt{\frac{f_i}{300kHz}} + (1 - a - b) \frac{f_i}{300kHz} \right) \quad [\text{dB}] \quad (27)$$

gdzie a , b są stałymi dobieranymi do danych empirycznych, to wówczas możemy wyrazić zależność przepływności B w systemie ADSL od tłumienia linii na częstotliwości 300 kHz, $Att(300 \text{ kHz})$ jako

$$B = Z - V \cdot Att(300kHz) \quad [\text{bit/s}] \quad (28)$$

gdzie Z i V są pewnymi stałymi, wyrażonymi przez

$$Z = \frac{m}{3} \left(\sum_i P_i - \sum_i N_i - i \cdot 10_{dB} \right) \quad (29)$$

$$V = \frac{m}{3} \sum_i \left(a + b \sqrt{\frac{f_i}{300kHz}} + (1 - a - b) \frac{f_i}{300kHz} \right) \quad (30)$$

Zastanówmy się teraz nad ograniczeniami w zastosowaniu zależności (28). Po pierwsze warunek $SNR_i \gg 10$ jest spełniony tylko dla dużych przepływności, a co za tym idzie przy małych przepływnościach przybliżenie (25) i dalej (28) nie ma sensu, gdyż formalne potraktowanie zależności (28) doprowadziłoby do uzyskania ujemnych przepływności przy dużych tłumieniach linii. Z kolei w praktycznym systemie wartość maksymalnej przepływności w kanale jest ograniczona wartościowością użytej modulacji i nie wzrasta nieograniczenie nawet przy $SNR_i \rightarrow \infty$. Zatem formuła (28) może być w praktyce stosowana jedynie dla oszacowania nachylenia spadku przepływności wraz ze wzrostem tłumienia. Dla uzyskania wartości liczbowych przeprowadzimy teraz stosowne obliczenia. Mamy

$$V = \frac{m}{3} \sum_i \left(a + b \sqrt{\frac{f_i}{300 \text{ kHz}}} + (1 - a - b) \frac{f_i}{300 \text{ kHz}} \right) = \frac{m}{3 \cdot \Delta f} \sum_i \left(a + b \sqrt{\frac{f_i}{300 \text{ kHz}}} + (1 - a - b) \frac{f_i}{300 \text{ kHz}} \right) \Delta f \approx \frac{m}{3 \cdot \Delta f} \int_{f_1}^{f_2} \left(a + b \sqrt{\frac{f}{300 \text{ kHz}}} + (1 - a - b) \frac{f}{300 \text{ kHz}} \right) df = \frac{m}{3 \cdot \Delta f} \left[a(f_2 - f_1) + b \frac{2}{3 \sqrt{300 \text{ kHz}}} (f_2^{3/2} - f_1^{3/2}) + (1 - a - b) \frac{1}{2 \cdot 300 \text{ kHz}} (f_2^2 - f_1^2) \right] \quad (31)$$

W powyższym wzorze $\Delta f = 4.3125 \text{ kHz}$ jest szerokością kanału w systemie ADSL, zaś f_1 i f_2 , odpowiednio górną i dolną częstotliwością graniczną. Wartości tych częstotliwości są następujące:

Dla kierunku upstream: $f_1 \approx 26 \text{ kHz}$, $f_2 \approx 138 \text{ kHz}$,

Dla kierunku downstream bez rozdziału częstotliwościowego: $f_1 \approx 26 \text{ kHz}$, $f_2 \approx 1104 \text{ kHz}$,

Dla kierunku downstream z rozdziałem częstotliwościowym: $f_1 \approx 138 \text{ kHz}$, $f_2 \approx 1104 \text{ kHz}$.

Wartości parametru V stosunkowo słabo zmieniają się przy zmianie parametrów a , b w modelu zależności tłumienia od częstotliwości (zmiany rzędu 20%). Dla kierunku downstream wartości $V \approx 0.4 \text{ Mbit/s/dB}$, zaś dla kierunku upstream $V \approx 24 \dots 29 \text{ kbit/s/dB}$. Te wielkości porównamy teraz z danymi doświadczalnymi. Możemy porównać funkcje (23) uzyskane empirycznie z zależnościami teoretycznymi (współczynnik nachylenia V). Mamy

$$\left. \frac{dy}{dx} \right|_{x=p} = -kB_{MAX}/4 \quad (32)$$

co daje dla kierunku downstream wartości rzędu $0.2 \dots 0.3 \text{ Mbit/s/dB}$, a więc nieco mniejsze niż wartość nachylenia $V (\approx 0.4 \text{ Mbit/s/dB})$ uzyskana na drodze teoretycznej. Jednak zważywszy na przybliżenia przyjęte przy wyprowadzeniu zależności teoretycznej można uznać zgodność uzyskanych wyników za niezłą. Nieco lepsza zgodność istnieje dla przepływności w kierunku upstream: dane pomiarowe 22 kbit/s/dB ; teoria $V = 24 \dots 29 \text{ kbit/s/dB}$.

5. PODSUMOWANIE

W pracy przedstawiono metodę kwalifikacji linii abonenckich do usług szerokopasmowych (ADSL) w oparciu o rezultaty sondowania linii modemami pracującymi w paśmie głosowym zgodnie z protokołem V.34. Istota tej metody polega na powiązaniu nachylenia charakterystyki przenoszenia toru (zdominowanego przez badaną pętlę abonencką) w paśmie głosowym z jego tłumieniem na częstotliwości 300 kHz. Charakterystyka ta jest mierzona przez modemy V.34 na etapie inicjalizacji połączenia. W pracy uzyskano odpowiednie zależności teoretyczne (punkt 2), które poprzez opracowany system pomiarowy wraz z oprogramowaniem zostały następnie zweryfikowane eksperymentalnie (punkt 3). Uzyskane wyniki pozwalają się spodziewać błędu w określeniu tłumienia linii abonenckiej na częstotliwości 300 kHz na poziomie nie większym niż ± 5 dB. Największy wpływ na ten błąd mają charakterystyki przenoszenia filtrów stosowanych na zakończeniach linii oraz w samych modemach. Mając obliczone tłumienie linii na częstotliwości 300 kHz można dokonać oszacowania spodziewanych do uzyskania przepływności w obydwu kierunkach transmisji w systemie ADSL (punkt 4.). Najbardziej sensowną metodą z punktu widzenia operatora wydaje się być zastosowanie przybliżenia funkcją sigmoidalną (23) z tak dobranymi parametrami, aby uzyskać zadane prawdopodobieństwo tego, iż rzeczywista przepływność linii nie będzie mniejsza od przepływności obliczonej. W pracy brak jest bezpośredniego powiązania pomiarów modemowych V. 34 z przepływnością systemów ADSL pracujących na tych samych pomierzonych liniach. Wykonujący tę pracę nie mieli możliwości przeprowadzenia takich badań, gdyż wymagałoby to uzyskania daleko idącej współpracy operatora telekomunikacyjnego. Ta współpraca ograniczyła się jedynie do umożliwienia porównania pomiarów modemowych i rzeczywistego tłumienia linii na częstotliwości 300 kHz, za co autorzy są wdzięczni Telekomunikacji Polskiej S.A.

Praca była możliwa dzięki grantowi KBN nr 4 T11D01525.

6. BIBLIOGRAFIA

1. IEC 61156-1, *Multicore and symmetrical pair/quad cables for digital communications*, Part 1: Generic specification, 2001
2. ITU-T G.996.1, *Test procedures for digital subscriber line transceivers*, 02/2001
3. ITU-T V.34, *A modem operating at data signalling rates of up to 33 600 bit/s for use on the general switched telephone network and on leased point-to-point 2-wire telephone-type circuits*, 02/1998
4. WT-062v8: *ADSL Interoperability Test Plan*, DSL Forum 2001
5. M. Budzyński, M. Ratkiewicz: *Testowanie systemów dostępowych xDSL dla potrzeb operatorów*, Sprawozdanie z pracowni dyplomowej, 2002
6. ANSI T1.413, *Network to Customer Installation Interfaces — Asymmetric Digital Subscriber Line (ADSL) Metallic Interfaces*, 1998
7. J. Siuzdak, A. Kalinowski: *Raport z pomiarów parametrów warstwy fizycznej modemów ADSL firm Siemens, Cisco, Ascom, Nokia, Alcatel, Paradyne, Ericsson, Lucent*, Instytut Telekomunikacji PW, Warszawa 2001 (raport dla TP S.A.)

8. M. Rydel: *Transmisja sygnałów w torach przewodowych*, Wydawnictwa Politechniki Warszawskiej, Warszawa 1980
9. ITU-T G.992.1, *Asymmetric digital subscriber line (ADSL) transceivers*, 1999
10. B. Farhang-Bouroujeny, M. Ding: *Design methods for time-domain equalizers in DMT transceivers*, IEEE Trans. Communications, vol. 49, no. 3, March 2001, pp.554-562
11. J. Cioffi, G. Dudevoir, M.V. Eyuboglu, G.D. Forney: *MMSE-DFEs and Coding*, Part I and II, IEEE Trans. Communications, vol. 43, no. 10, October 1995, pp. 2582-2604.
12. TU-T, Study Group 15, Question 4/15, G.vdsl: *Construction of modulated signals from filter-bank elements and equivalence of line codes*, August 1999
13. N.S. Alagha: *Modulation, pre-equalization and pulse shaping for PCM voiceband channels*, Ph.D. thesis, McGill University, Montreal Canada, December 2001
14. Ł. Mikocki, A. Dulak: Raport nr CBR-TKR/L/063/11/2004, Laboratorium Pomiarowe Telekomunikacyjnej Izby Pomiarowej w Krakowie
15. A. Majewski, P. Czachorowski: *Praktyczne zastosowanie techniki ADSL*, Praca dyplomowa, Instytut Radioelektroniki, Politechnika Warszawska, 2004/2005
16. J.L. Lechleider, S. Narain, C.H. Woloszynski: *Method and system for estimating the ability of a subscriber loop to support broadband services*, United States Patent no. 6091713, 18 July 2000

J. SIUZDAK, T. CZARNECKI

SUBSCRIBER LINE QUALIFICATION FOR ADSL BY MEANS OF V.34 MODEMS

S u m m a r y

The paper presents the method of subscriber line qualification for ADSL service based on the results of line probing by V.34 voice band modems. The line frequency response is measured in the voice band by V.34 modems during the connection initialization stage. Then this response is analyzed, and its slope calculated. The calculated slope is related to the line attenuation at 300 kHz. The respective relations are obtained theoretically, and presented in the paper. The method was practically implemented, and the measurement results are also shown. The results indicate that the estimate of the line attenuation at 300 kHz obtained in this way does not differ from the real value by more than 5 dB. Given the line attenuation it is possible to assess the down and up bit rates of ADSL operating over the same line. To this end the sigmoidal function of the attenuation was used with the parameters chosen to minimize the mean square error of approximation.

Keywords: data transmission, access networks, line qualification, broadband services, ADSL, V.34, subscriber loop, attenuation

I

C

J

bocze
skutecz
między
przewo
peratu
2]. R
badan

Estymacja parametrów modelu termicznego elementów półprzewodnikowych

KRZYSZTOF GÓRECKI, JANUSZ ZARĘBSKI

*Katedra Elektroniki Morskiej, Akademia Morska w Gdyni,
ul. Morska 83, 81-225 Gdynia
e-mail: gorecki@am.gdynia.pl, zarebski@am.gdynia.pl*

Otrzymano 2005.09.01

Autoryzowano 2005.11.03

W pracy przedstawiono opracowaną przez autorów metodę estymacji parametrów modelu termicznego elementów półprzewodnikowych i układów scalonych oraz program komputerowy MASTER2 realizujący tę metodę. Program MASTER2 steruje procesem pomiaru przejściowej impedancji termicznej $Z(t)$ badanych elementów półprzewodnikowych i na podstawie uzyskanego przebiegu $Z(t)$ wyznacza wartości parametrów skupionego modelu termicznego - rezystancji termicznej, termicznych stałych czasowych i odpowiadających im współczynników wagowych, a także wartości elementów biernych występujących w analogu elektrycznym modelu termicznego sformułowanego w postaci sieci Cauera lub Fostera. Poprawność opracowanej metody i programu zweryfikowano doświadczalnie przez porównanie zmierzonych i obliczonych przebiegów przejściowej impedancji termicznej $Z(t)$ tranzystora mocy MOS.

Słowa kluczowe: przejściowa impedancja termiczna, model termiczny, pomiary termiczne

1. WPROWADZENIE

Jednym z istotnych czynników wpływających na charakterystyki i parametry robocze przyrządów półprzewodnikowych oraz ich niezawodność jest temperatura. Na skutek zjawisk termicznych — samonagrzewania i wzajemnych sprzężeń termicznych między elementami półprzewodnikowymi umieszczonymi na wspólnym podłożu półprzewodnikowym, na wspólnym radiatorze lub na wspólnej płytce drukowanej, temperatura wnętrza elementów może być znacznie wyższa od temperatury otoczenia [1, 2]. Różnica między wymienionymi temperaturami zależy od warunków chłodzenia badanych elementów.

W skupionym modelu termicznym elementu półprzewodnikowego zależność temperatury jego wnętrza T_j od wydzielanej w nim mocy można wyrazić za pomocą całki spłotu o postaci

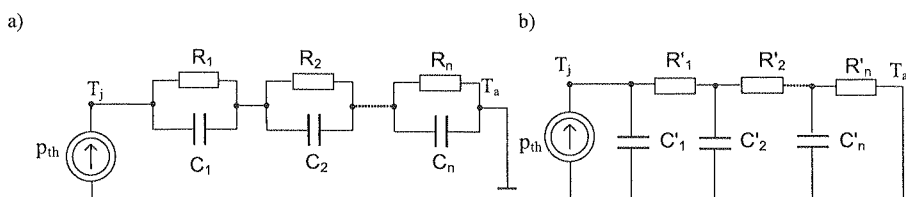
$$T_j = T_a + \int_0^t Z'(t-v) \cdot p(v) \cdot dv \quad (1)$$

gdzie T_a oznacza temperaturę otoczenia, $p(v)$ moc czynną wydzielaną w rozważanym elemencie półprzewodnikowym, natomiast $Z'(t)$ jest czasową pochodną przejściowej impedancji termicznej $Z(t)$ tego elementu, opisywaną zwykle wzorem [1, 2, 3, 4]

$$Z(t) = R_{th} \cdot \left[1 - \sum_{i=1}^N a_i \cdot \exp\left(-\frac{t}{\tau_{thi}}\right) \right] \quad (2)$$

gdzie R_{th} oznacza rezystancję termiczną, a_i to współczynniki wagowe odpowiadające poszczególnym termicznym stałym czasowym τ_{thi} , natomiast N jest liczbą tych stałych czasowych.

W wielu programach komputerowych, umożliwiających analizę elektrotermiczną układów elektronicznych, np. APLAC [5], nie można bezpośrednio stosować wzoru (1), lecz wykorzystuje się ekwiwalent obwodowy modelu termicznego elementów półprzewodnikowych. Ekwiwalentem obwodowym równań (1) i (2) są sieci Foster'a (rys. 1a) lub Cauera (rys. 1b). Na rys. 1 źródło prądowe p_{th} reprezentuje moc cieplną wydzielaną w elemencie, napięcie w węźle T_j odpowiada nadwyżce temperatury wnętrza elementu ponad temperaturę otoczenia, natomiast sieć RC reprezentuje przejściową impedancję termiczną rozważanego elementu. Jak wynika m.in. z prac [1, 3], obie rozważane



Rys. 1. Model termiczny w postaci sieci Foster'a (a) oraz sieci Cauera (b)

Fig. 1. The network thermal models of the Foster (a) and Cauer (b) form

sieci są w pełni równoważne z punktu widzenia zacisku T_j (rys. 1), przy czym sieć Foster'a nie ma bezpośredniej interpretacji fizycznej, natomiast struktura sieci Cauera wynika wprost z dyskretyzacji jednowymiarowego równania przewodnictwa ciepła, a napięcia w poszczególnych węzłach tej sieci odpowiadają temperaturom elementów konstrukcyjnych badanego elementu półprzewodnikowego, np. podstawki montażowej czy obudowy. W literaturze znane są przykłady zastosowania do wyznaczania temperatury wnętrza elementów półprzewodnikowych obu rozważanych struktur modelu

termicznego, zarówno sieci Fostera [1, 3, 6], jak i Cauera [3, 7]. Istotną zaletą sieci Fostera jest możliwość uzyskania wartości elementów tej sieci bezpośrednio na podstawie wartości parametrów modelu termicznego występujących we wzorze (2). Rezystancja R_i oraz pojemność C_i dane są wówczas wzorami

$$R_i = a_i \cdot R_{th} \quad (3)$$

$$C_i = \tau_{thi}/R_i \quad (4)$$

Wyznaczenie wartości parametrów modelu termicznego elementu półprzewodnikowego w każdej z wymienionych postaci wymaga więc wykonania najpierw pomiaru przejściowej impedancji termicznej tego elementu, a następnie odpowiednich obliczeń przy wykorzystaniu zmierzonego przebiegu $Z(t)$ badanego elementu.

Pomiary $Z(t)$ są realizowane przy wykorzystaniu specjalizowanych przyrządów lub systemów pomiarowych [2, 8–13]. Większość z nich [2, 8, 9, 13] pozwala na pomiar wyłącznie rezystancji termicznej jednego typu elementów półprzewodnikowych, np. diody $p-n$ [2, 10, 11, 12], tranzystora bipolarnego [8, 13], czy też tranzystora Darlingtona [9]. Z kolei, opracowany i uruchomiony przez autorów prototyp systemu pomiarowego [14], pozwala na pomiary przejściowej impedancji termicznej zarówno diod $p-n$, tranzystorów bipolarnych i unipolarnych, jak również tranzystorów Darlingtona, przy wykorzystaniu różnych metod pomiarowych, opisanych m.in. w pracach [15–17].

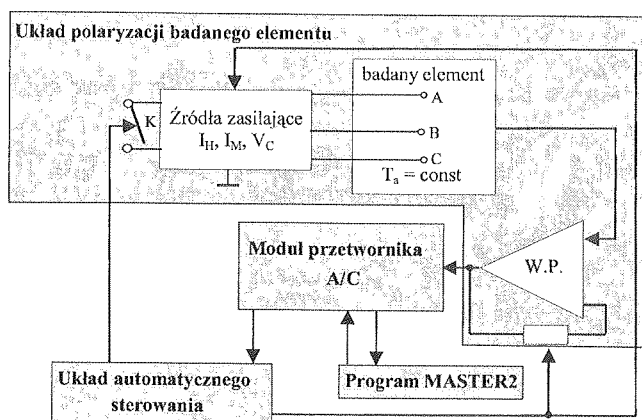
Literaturowe metody wyznaczania wartości parametrów R_{th} , a_i , τ_{thi} modelu termicznego, na podstawie wyników pomiarów $Z(t)$, mają szereg wad, a niekiedy są uciążliwe w realizacji. Przykładowo, algorytm z pracy [4] wymaga złożonych obliczeń przy wykorzystaniu algorytmu rozplotowego, który nie zawsze jest zbieżny, natomiast algorytm heurystyczny, zaproponowany w pracy [18] wymaga założenia a priori wartości termicznych stałych czasowych. Z kolei, w pracy [6] zaproponowano algorytm wyznaczania wartości elementów występujących w analogu obwodowym modelu termicznego w postaci sieci Fostera. Algorytm ten jest uciążliwy w realizacji, ponieważ wymaga między innymi wyznaczenia tzw. „widma termicznych stałych czasowych” przy wykorzystaniu algorytmów rozplotowych. W pracy [3] opisano metodę wyznaczania wartości elementów sieci Cauera w oparciu o znany opis zespolonej impedancji termicznej $Z(s)$, przy czym nie podano szczegółów dotyczących sposobu uzyskania opisu $Z(s)$ na podstawie zmierzonego przebiegu $Z(t)$. Z kolei, opracowany wcześniej przez autorów interaktywny algorytm wyznaczania wartości parametrów modelu termicznego danego wzorem (2) wymaga od użytkownika żmudnej analizy zmierzonego przebiegu $Z(t)$ i wyznaczenia przedziałów liniowości pewnej funkcji, której argumentem jest $Z(t)$ [19].

Celem pracy było opracowanie metody estymacji parametrów modelu termicznego elementów półprzewodnikowych, jej implementacja w autorskim programie komputerowym MASTER2 oraz weryfikacja jej poprawności. Program MASTER2 steruje procesem pomiaru przebiegu $Z(t)$ w systemie pomiarowym [14], a następnie wyznacza także wartości parametrów modelu termicznego danego wzorem (2), jak również

wartości elementów RC w analogu obwodowym modelu termicznego sformułowanego zarówno w postaci sieci Cauera, jak i Fostera.

W dalszej części niniejszej pracy przedstawiono strukturę systemu pomiarowego do wyznaczania impedancji termicznej elementów półprzewodnikowych, sterowanego przez program MASTER2, opis sieci działań tego programu, zaprezentowano metodę estymacji parametrów, występujących we wzorze (2) oraz wartości elementów RC sieci Cauera i Fostera, a także przedstawiono wyniki weryfikacji poprawności opracowanej metody oraz działania programu MASTER2.

2. SYSTEM POMIAROWY



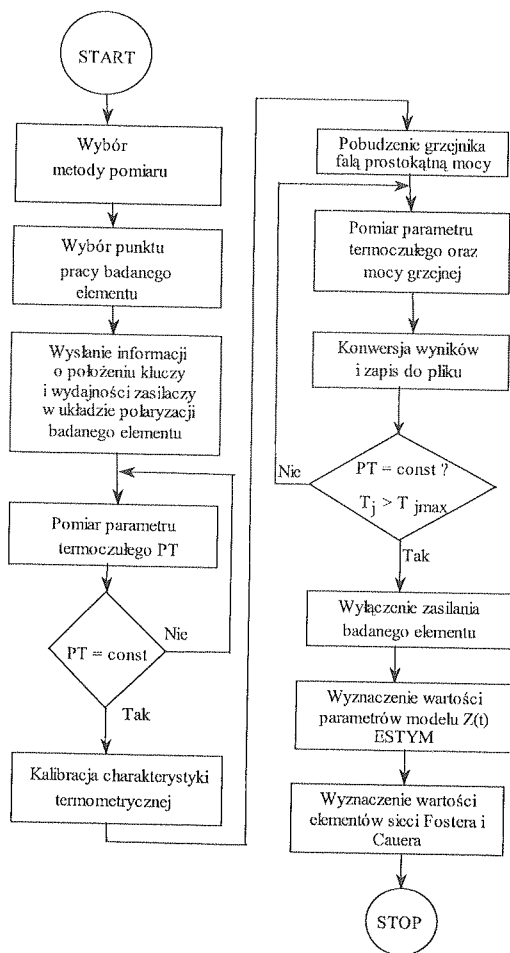
Rys. 2. Schemat blokowy mikrokomputerowego systemu pomiarowego

Fig. 2. The block diagram of the microcomputer measuring set

W celu zmierzenia przebiegu $Z(t)$ wykorzystywany jest system pomiarowy, którego schemat blokowy przedstawiono na rys. 2 [14]. Składa się on z czterech zasadniczych bloków: komputera PC, układu sterującego, układu polaryzacji badanego elementu oraz wzmacniaczy pomiarowych. Komputer zawiera specjalizowany moduł przetwornika analogowo-cyfrowego, umożliwiający pomiar odpowiednich napięć i prądów zaciskowych badanego elementu oraz układy wejść i wyjść cyfrowych, zapewniających połączenie komputera z układem sterowania, wytwarzającym sygnały sterujące wydajnościami zasilaczy prądowych i napięciowych oraz przełącznikami w układzie polaryzacji badanego elementu, wzmacnieniami wzmacniaczy pomiarowych oraz termostatem, zapewniającym stałą wartość temperatury otoczenia w czasie pomiaru. Wymienione bloki funkcjonalne szczegółowo opisano w pracy [14]. System pomiarowy jest kontrolowany przez, opisany w rozdziale 3, specjalizowany program MASTER2, napisany w języku Borland DELPHI.

3. PROGRAM MASTER2

Działanie systemu pomiarowego, przedstawionego w poprzednim rozdziale, nadzoruje program MASTER2, którego sieć działań przedstawiono na rys. 3. Przed rozpoczęciem pomiarów przejściowej impedancji termicznej użytkownik wprowadza do programu MASTER2 niezbędne informacje o procesie pomiarowym, m. in.: nazwę pliku wynikowego, wybraną metodę pomiaru, wartość napięcia i prądu wydzielanego w badanym elemencie w czasie jego nagrzewania. Program realizuje proces pomiarowy poprzez wysyłanie przez układ wyjść TTL, znajdujący się w module przetwornika A/C, odpowiednich ciągów binarnych do układu automatycznego sterowania oraz odczytywanie ciągów bitów uzyskiwanych z wyjścia przetwornika A/C.



Rys. 3. Sieć działań programu MASTER2

Fig. 3. The diagram of the MASTER2

W pierwszym etapie pomiaru przez badany element płynie prąd pomiarowy o małej wartości i mierzona jest wartość parametru termoczułego PT przy temperaturze wnętrza elementu równej temperaturze otoczenia. Po ustaleniu się wartości parametru PT wyznaczane jest nachylenie F charakterystyki termometrycznej $PT(T)$ przy wykorzystaniu kalibracji jednopunktowej [20]. W charakterze parametru termoczułego wykorzystuje się napięcie na spolaryzowanym w kierunku przewodzenia złączu p-n, przy przepływie przez nie prądu o ustalonej, małej wartości.

W drugim etapie pomiaru $Z(t)$ badany element pobudzany jest falą prostokątną mocy o wypełnieniu d bliskim jedności ($d = 0,999$), co zapewnia praktycznie pobudzenie badanego elementu uskokiem mocy. Przy wysokim poziomie wydzielanej mocy mierzone są wartości zaciskowych prądów i napięć elementu, niezbędne do określenia wartości wydzielanej w nim mocy, natomiast przy niskim poziomie mocy mierzone są wartości PT, niezbędne do wyznaczenia wartości temperatury wnętrza badanego elementu. Zmierzone wartości zapisywane są do pliku wynikowego. Pomiary trwają do chwili ustalenia się wartości parametru PT lub przekroczenia dopuszczalnej wartości temperatury wnętrza T_{max} . Przy spełnieniu drugiego z wymienionych warunków następuje automatyczne wyłączenie zasilania badanego elementu, co zabezpiecza ten element przed ewentualnym zniszczeniem.

Na podstawie zmierzonego przebiegu czasowego $PT(t)$ wyznaczany jest przebieg przejściowej impedancji termicznej ze wzoru

$$Z(t) = \frac{PT(t) - PT(t=0)}{P_0} \cdot F^{-1} \quad (5)$$

gdzie P_0 jest wysokością uskoku mocy wydzielanej w badanym elemencie.

Na podstawie uzyskanego przebiegu $Z(t)$ wyznaczane są wartości parametrów, występujących w równaniu (2) przy wykorzystaniu opisanej w kolejnym rozdziale procedury ESTYM, a następnie wyznaczane są wartości elementów występujących w sieciach Cauera i Fostera.

4. ESTYMACJA PARAMETRÓW MODELU $Z(t)$

Danymi wejściowymi dla procedury ESTYM, stanowiącej część programu MASTER2, są zmierzone przebiegi $Z(t)$. Procedura ESTYM wykorzystuje opracowany przez autorów algorytm estymacji parametrów modelu $Z(t)$, opisany w pracach [1, 19, 21]. W procedurze ESTYM wartość R_{th} odpowiada wartości $Z(t)$ w stanie ustalonym. W celu wyznaczenia wartości parametrów a_i oraz τ_{thi} zdefiniowano funkcję $y_i(t)$ o postaci

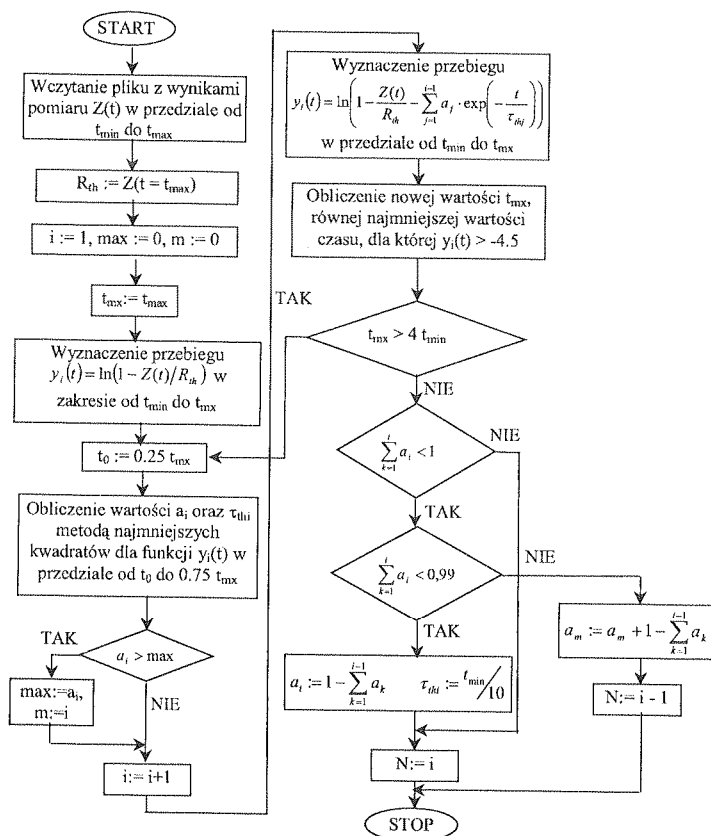
$$y_i(t) = \ln \left(1 - \frac{Z(t)}{R_{th}} - \sum_{j=1}^{i-1} a_j \cdot \exp \left(-\frac{t}{\tau_{thj}} \right) \right) \quad (6)$$

Ponieważ termiczne stałe czasowe znacznie różnią się od siebie, dla dużych wartości czasu t , o przebiegu $Z(t)$ decyduje tylko najdłuższa termiczna stała czasowa τ_{thl} ,

a czynniki eksponencjalne we wzorze (2), związane z krótszymi stałymi czasowymi są pomijalnie małe ($t \gg \tau_{thi}$). Wówczas zależność (6) redukuje się do zależności liniowej o postaci

$$y_i(t) = -\frac{t}{\tau_{thi}} + \ln(a_i) \quad (7)$$

Wyznaczenie wartości parametrów τ_{thi} oraz a_i w zależności (7) wymaga zastosowania metody najmniejszych kwadratów, przy czym do aproksymacji należy wykorzystać tylko współrzędne punktów, leżących w zakresie liniowości zależności (6). Obliczenia są realizowane sekwencyjnie, począwszy od najdłuższej termicznej stałej czasowej, ku coraz krótszym termicznym stałym czasowym. Przy wyznaczaniu parametrów τ_{thi} oraz a_i wykorzystywane są wyznaczone w poprzednich krokach wartości tych parametrów związanych z dłuższymi, od obliczanej, termicznymi stałymi czasowymi. Sieć działań procedury ESTYM, realizującej powyżej przedstawiony algorytm, przedstawiono na rys.4.



Rys. 4. Sieć działań procedury ESTYM

Fig. 4. The diagram of the ESTYM

Wyznaczanie wartości parametrów modelu przejściowej impedancji termicznej rozpoczyna się od określenia wartości rezystancji termicznej R_{th} przez uśrednienie zmierzonych wartości $Z(t)$ w stanie ustalonym. Przyjęto, że stan taki występuje dla ostatnich 100 sekund mierzonego przebiegu $Z(t)$.

Następnie wyznacza się wartości τ_{th1} oraz a_1 , odpowiadające najdłuższej termicznej stałej czasowej na podstawie współrzędnych punktów z przedziału liniowości charakterystyki $y_1(t)$. Ze względu na duże różnice między wartościami kolejnych termicznych stałych czasowych przyjęto, że przedział ten obejmuje czas od $t_o = 25\%t_{mx}$ do $75\% t_{mx}$, przy czym dla $i = 1$ czas t_{mx} równy jest czasowi zakończenia pomiarów t_{mx} .

Po wyznaczeniu wartości a_1 oraz τ_{th1} wyznaczana jest zależność $y_2(t)$ ze wzoru (6), a nowa wartość czasu t_{mx} jest najmniejszą liczbą spełniającą warunki

$$y_2(t_{mx}) < -4,5 \quad (8)$$

$$t_{mx} < \frac{\tau_{th(i-1)}}{4} \quad (9)$$

Warunek (8) wynika z założenia, że dowolny współczynnik a_i jest większy od 0,01, natomiast warunek (9) — z założenia, iż $\tau_{th(i-1)}/\tau_{thi} \gg 1$.

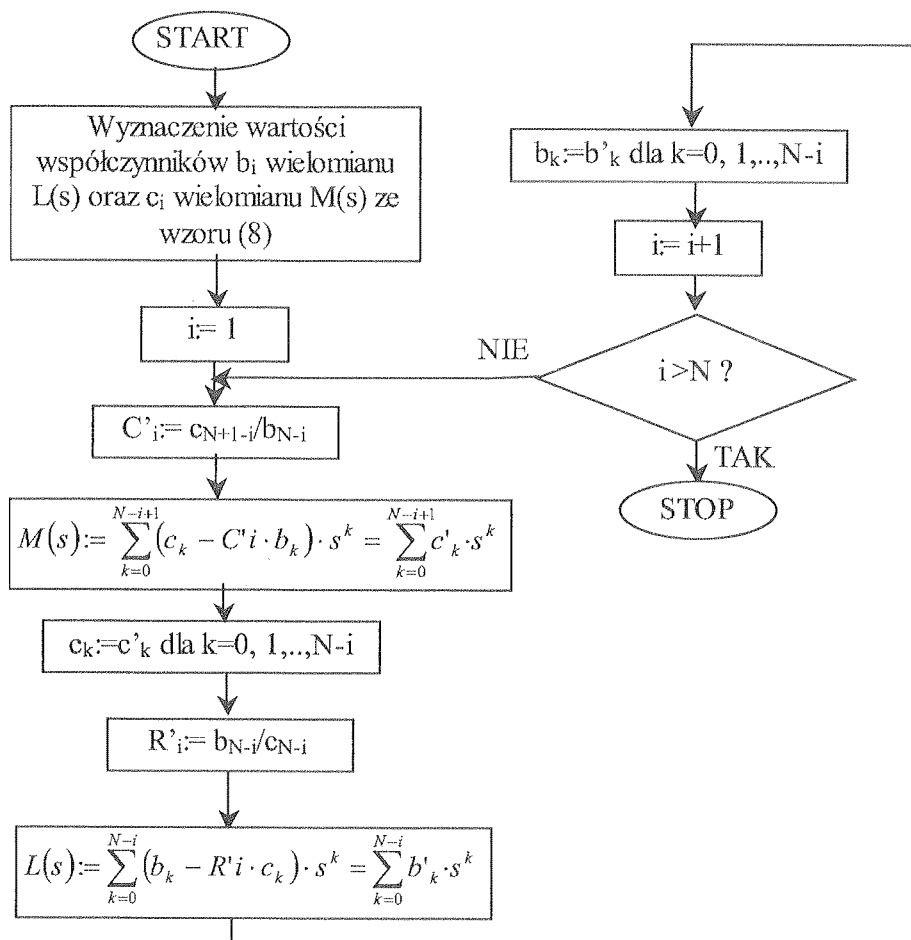
Po obliczeniu parametrów a_2 oraz τ_{th2} wyznaczana jest nowa wartość czasu t_{mx} . Jeżeli obliczona wartość czasu $t_{mx} > 4t_{min}$, gdzie t_{min} jest czasową współrzędną pewnego zmierzonego punktu w przebiegu $Z(t)$, to wyznaczane są parametry a_i oraz τ_{thi} dla kolejnej wartości indeksu i dla przedziału czasu od $25\%t_{mx}$ do $75\%t_{mx}$ przy nowej wartości t_{mx} spełniającej warunki (8) i (9). Proces ten powtarzany jest iteracyjnie tak długo, aż ostatnia wyznaczona wartość czasu t_{mx} będzie mniejsza od $4t_{min}$. Przed zakończeniem procedury wyliczana jest suma wyznaczonych dotychczas wartości a_i . Jeżeli suma ta jest mniejsza od 1, ale większa od 0,99, to przyjmuje się, że wystąpił błąd zaokrągleń. Wówczas, w celu zapewnienia zerowej wartości $Z(t)$ dla czasu $t = 0$, zwiększana jest wartość największego z wyznaczonych współczynników a_i (na rys.4 oznaczonego symbolem a_m) w taki sposób, aby suma tych współczynników była równa 1. Z kolei, mniejsza od 0,99 wartość rozważanej sumy oznacza, że badany element posiada niemierzalną za pomocą systemu pomiarowego najkrótszą termiczną

stałą czasową. Wówczas wprowadza się kolejny współczynnik $a_i = 1 - \sum_{k=1}^{i-1} a_k$, z którym skojarzona jest termiczna stała czasowa o arbitralnie wybranej wartości $\tau_{thi} = t_{min}/10$.

Następnie, uzyskane przy wykorzystaniu procedury ESTYM, wartości parametrów R_{th} , a_i oraz τ_{thi} są wykorzystywane do obliczenia wartości elementów biernych, występujących w analogach obwodowych modelu termicznego. Wartości elementów sieci Fostera są odpowiednio obliczane ze wzorów (3) i (4), natomiast wyznaczenie wartości elementów w sieci Cauera jest bardziej złożone i wymaga:

a) Wy
wz

gdz
b) Na
ko
tej



Rys. 5. Wyznaczenie wartości elementów RC sieci Cauera

Fig. 5. The diagram of computing of RC elements of the Cauer network

- a) Wyznaczenia operatorowej impedancji termicznej $Z(s)$ badanego elementu danej wzorem

$$Z(s) = R_{th} \cdot \frac{\sum_{i=1}^N \left(\frac{a_i}{\tau_{thi}} \cdot \prod_{j=1, j \neq i}^N \left(s + \frac{1}{\tau_{thj}} \right) \right)}{\prod_{j=1}^N \left(s + \frac{1}{\tau_{thj}} \right)} = \frac{L(s)}{M(s)} = \frac{\sum_{i=0}^{N-1} b_i \cdot s^i}{\sum_{i=0}^N c_i \cdot s^i} \quad (10)$$

gdzie N oznacza liczbę termicznych stałych czasowych w modelu $Z(t)$.

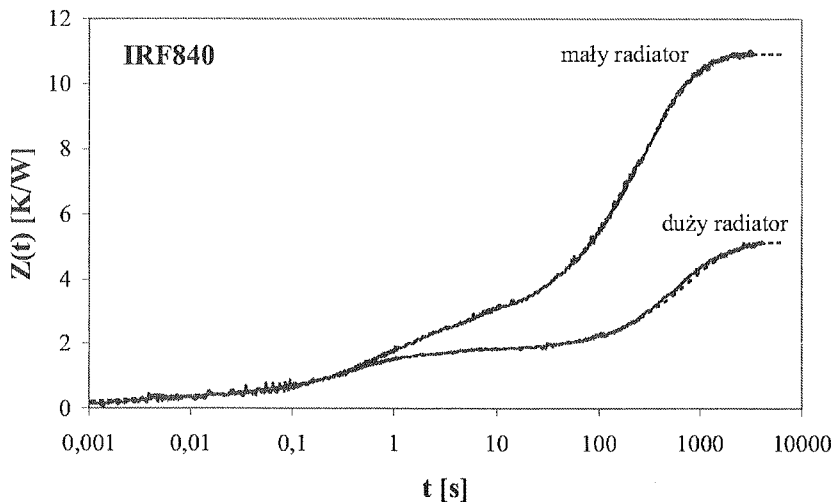
- b) Naprzemiennego dzielenia wielomianów $L(s)$ i $M(s)$, przy czym dzielone są tylko wyrazy występujące przy najwyższych stopniach zmiennej s . Sposób realizacji tej operacji zilustrowano na rys. 5. W przypadku, gdy dzielone są wielomiany o

równych stopniach uzyskuje się wartość rezystancji R'_i , natomiast gdy stopień licznika jest większy od stopnia mianownika o 1, to uzyskuje się wartość pojemności C'_i .

5. WYNIKI BADAŃ

W celu weryfikacji poprawności opracowanej metody estymacji i programu MASTER2 wykonano serię pomiarów przejściowej impedancji termicznej wybranych elementów półprzewodnikowych oraz wyznaczono wartości parametrów modelu termicznego i wartości elementów RC, występujących w sieciach Cauera i Fostera, odpowiadających tym parametrom.

Na rys.6 przedstawiono wyniki badań tranzystora MOS typu IRF840 umieszczonego kolejno na dwóch radiatorach wykonanych z kształtki A-4240 — pierwszy o długości 60 mm (duży radiator) oraz drugi o długości 18 mm (mały radiator). Na rysunku tym linie ciągłe oznaczają zmierzone przebiegi $Z(t)$. Z kolei, ze względu na to, że przebiegi $Z(t)$ uzyskane z obliczeń przy wykorzystaniu wzoru (2) oraz przy użyciu obu analogów obwodowych są nierozróżnialne, na rys.6 oznaczono je wspólnie za pomocą linii kreskowych, także praktycznie pokrywających się z danymi pomiarowymi. Wartości parametrów modelu $Z(t)$ oraz elementów sieci Cauera i Fostera podano w tabeli 1 dla tranzystora umieszczonego na małym radiatorze oraz w tabeli 2 dla tranzystora na dużym radiatorze.



Rys. 6. Zmierzone (linie ciągłe) i obliczone (linie kreskowe) przebiegi $Z(t)$ tranzystora IRF840

Fig. 6. The simulated (dashed lines) and measured (solid lines) values of $Z(t)$ of the VDMOS transistor IRF840

licz-
ności

MA-
ch ele-
micz-
powia-

SZCZO-
wszy o
(r). Na
ędu na
z przy
pólnie
pomia-
Fostera
Tabeli 2

Jak widać, krzywe obliczone zarówno za pomocą wzoru (2), jak również otrzymane z analogów obwodowych modelu termicznego z zadawalającą dokładnością aproksymują zmierzony przebieg przejściowej impedancji termicznej, co świadczy o poprawności działania programu. Różnice między wynikami pomiarów i obliczeń nie przekraczają 0,1 K/W. Dokładny opis przebiegów $Z(t)$ wymaga użycia 8 termicznych stałych czasowych. Czas potrzebny na uzyskanie stanu ustalonego silnie zależy od wielkości radiatora i w rozpatrywanych przypadkach zmienia się od 500 s dla tranzystora na małym radiatorze do 3000 s dla tranzystora na dużym radiatorze. Obydwa rozważane przebiegi przejściowej impedancji termicznej są praktycznie identyczne dla czasów $t < 0,1$ s, co może oznaczać, iż w tym przedziale czasu o przebiegu $Z(t)$ decydują właściwości samego elementu, a nie radiatora.

Tabela 1

Wartości parametrów modelu termicznego tranzystora IRF840 z małym radiatorem

Values of the thermal model of IRF840 transistor situated on the heat-sink of the little area

Model $Z(t) - R_{th} = 10.92 \text{ K/W}$		Sieć Fostera		Sieć Cauera	
a_i	$\tau_{thi} [\text{s}]$	$R_i [\Omega]$	$C_i [\text{F}]$	$R'_i [\Omega]$	$C'_i [\text{F}]$
0.392	467.7	4.28	109.2	0.215	1.14×10^{-2}
0.298	176.8	3.25	54.33	0.925	2.49×10^{-1}
0.035	110.3	0.382	288.6	0.942	4.33×10^{-1}
0.033	18.1	0.360	50.23	0.861	4.26
0.079	3.411	0.863	3.954	2.11	19.03
0.112	0.643	1.223	0.526	4.26	17.3
0.02	0.128	0.218	0.584	1.16	207.1
0.018	0.00235	0.197	0.01196	0.308	586.4

Tabela 2

Wartości parametrów modelu termicznego tranzystora IRF840 z dużym radiatorem

Values of the thermal model of IRF840 transistor situated on the heat-sink of the large area

Model $Z(t) - R_{th} = 5.16 \text{ K/W}$		Sieć Fostera		Sieć Cauera	
a_i	$\tau_{thi} [\text{s}]$	$R_i [\Omega]$	$C_i [\text{F}]$	$R'_i [\Omega]$	$C'_i [\text{F}]$
0.404	1007.4	2.084	483.4	6.435×10^{-2}	6.62×10^{-4}
0.247	492	1.274	386.1	0.308	1.04×10^{-2}
0.06	2.756	0.309	8.906	0.381	1.57×10^{-1}
0.197	0.383	1.016	0.376	0.829	2.82×10^{-1}
0.028	0.0377	0.144	0.261	0.247	11.26
0.053	0.00317	0.273	0.0116	2.938	205.1
0.011	0.00004	0.057	0.000705	0.389	1998

840
OS

6. PODSUMOWANIE

W pracy przedstawiono metodę estymacji parametrów termicznych elementów półprzewodnikowych oraz program komputerowy MASTER2, realizujący tę metodę. Program ten umożliwia wyznaczenie wartości parametrów występujących w modelu przejściowej impedancji termicznej, danej wzorem (2), oraz wyznaczenie wartości elementów RC występujących w analogach elektrycznych modelu termicznego tego elementu, a także steruje procesem pomiaru przejściowej impedancji termicznej przy wykorzystaniu mikrokomputerowego systemu pomiarowego z pracy [14]. Proces estymacji parametrów modelu termicznego realizowany jest w pełni automatycznie, a przedstawione wyniki badań potwierdzają poprawność działania opracowanego programu.

7. PODZIĘKOWANIA

Pracę zrealizowano w ramach projektu nr 3 T10A 032 26 finansowanego przez Komitet Badań Naukowych w latach 2004-2005.

8. BIBLIOGRAFIA

1. J. Zarębski: *Modelowanie, symulacja i pomiary przebiegów elektrotermicznych w elementach półprzewodnikowych i układach elektronicznych*, Prace Naukowe Wyższej Szkoły Morskiej w Gdyni, Gdynia, 1996.
2. A. Nowakowski: *Badanie procesów termicznych w przyrządach półprzewodnikowych*, Zesz. Nauk. Polit. Gdańskiej, Elektronika LX, nr 389, 1984.
3. P.E. Bagnoli, C. Casarosa, M. Ciampi, E. Dallago: *Thermal Resistance Analysis by Induced Transient (TRAIT) Method for Power Electronic Devices Thermal Characterization — Part I: Fundamentals and Theory*, IEEE Transactions on Power Electronics, Vol. 13, No. 6, 1998, p. 1208.
4. V. Székely: *Identification of RC Networks by Deconvolution: Chances and Limits*, IEEE Transactions on Circuits and Systems-I, Vol. 45, No. 3, 1998, p. 244.
5. T. Veijola, M. Valtonen: *Combined Electrical and Thermal Circuit Simulation Using APLAC. Part A: Implementation and Basic Components Models*, Helsinki University of Technology, CT-25, 1995.
6. V. Székely: *A New Evaluation Method of Thermal Transient Measurement Results*, Microelectronic Journal, Vol. 28, 1997, No. 3, 1997, p. 277.
7. F.N. Masana: *Pseudo-3D Dynamic Thermal Model for Semiconductor Packages*, 7-th International Conference Mixed Design of Integrated Circuits and Systems MIXDES 2000, Gdynia, 2000, p. 357.
8. J. Glinianowicz, F. Christiaens, E. Beyne, Z. Staszak: *Measurement Circuits for Transient Thermal Characterisation on Power Transistors and Transformers*, XVII-th National Conference Circuit Theory and Electronic Networks, Wrocław - Polanica Zdrój, 1994, p. 359.
9. Z. Lisik: *Pomiar rezystancji cieplnej bipolarnych tranzystorów mocy Darlingtona*, Kwartalnik Elektroniki i Telekomunikacji, 1994, 40, z. 3, s. 369.
10. V. Székely, M. Rencz, B. Courtois: *Thermal investigations of ICs and Microstructures*, Sensors and Actuators; A Physical, Vol. A 71, No. 1-2, 1998, p. 1.

11. V. S. ...
Cool
12. V. S. ...
Mic
13. W. ...
Meth
14. K. ...
elem
s. 37
15. F.F. ...
stand
16. J. Z. ...
with
Mana
17. K. ...
Darlt
18. W. J. ...
jego
Gdań
19. J. Z. ...
Simul
1993,
20. J. Z. ...
nikow
21. J. Z. ...
EDPE

This paper
devices. A
controls th
values of Z
electrical an
The proced
the authors'
of the parame
using the m
are compute
The consider
differ from e
linear functi
of the rema

11. V. Székely: *Thermal Testing and Control by Means of Built-in Temperature Sensors*, Electronics Cooling, Vol. 4, No. 3, 1998, p. 36.
12. V. Székely, M. Rencz, B. Courtois: *Thermal Investigations of IC's and Microstructures*, Microelectronics Journal, Vol. 28, No 3, 1997, p. 205.
13. W. Janke, A. Jurewicz: *Fast Measurements of Bipolar Transistor Thermal Resistance by D.C. Method*, Workshop Microtechnology and Thermal Problems in Electronics, Zakopane, 1998, p. 34.
14. K. Górecki, J. Zarębski: *System mikrokomputerowy do pomiaru parametrów termicznych elementów półprzewodnikowych i układów scalonych*, Metrologia i Systemy Pomiarowe, Nr 4, 2001, s. 379.
15. F.F. Oettinger, D.L. Blackburn: *Semiconductor Measurement Technology: Thermal Resistance Measurements*, U. S. Department of Commerce, NIST/SP-400/86, 1990.
16. J. Zarębski, K. Górecki: *A Method of the BJT Transient Thermal Impedance Measurement with the Double Junction Calibration*, Proc. 11-th IEEE Semiconductor Thermal Measurement and Management Symposium (SEMI-THERM), San Jose (USA), 1995, p.80.
17. K. Górecki, J. Zarębski: *Impulsowe metody pomiaru parametrów termicznych tranzystora Darlingтона mocy*, Metrologia i Systemy Pomiarowe, t. VII, z. 3, 2000, s. 287.
18. W. Janke, W. Pietrenko: *Model analityczny impedancji termicznej i algorytm identyfikacji jego współczynników*, Raport Badawczy Nr 82/95, Katedra Elektroniki Ciała Stałego, Politechnika Gdańska, Gdańsk, 1995.
19. J. Zarębski, K. Górecki: *Thermal Transients in Bipolar Transistors; Measurements and Simulations*, Int. Conf. on Information, Systems Methods Applied to Engineering Problems, Malta 1993, Vol. 2, p. 111.
20. J. Zarębski, K. Górecki: *Badanie charakterystyk termometrycznych elementów półprzewodnikowych ze złączem p-n*, Metrologia i Systemy Pomiarowe, Nr 4, 2001, s. 397.
21. J. Zarębski, K. Górecki: *Investigations of Power MOSFETs Transient Thermal Impedance*, EDPE 2005, Dubrownik, 2005, p. 84, E-05-75.

K. GÓRECKI, J. ZARĘBSKI

PARAMETERS ESTIMATION OF THERMAL MODEL OF SEMICONDUCTOR DEVICES

S u m m a r y

This paper deals with the problem of the estimation of the thermal model parameters of semiconductor devices. A new computer tool — MASTER2 for estimation of these parameters is presented. MASTER2 controls the method of measuring the transient thermal impedance $Z(t)$ as well as it estimates both the values of $Z(t)$ parameters: R_{th} , a_i , τ_{thi} ($i = 1, 2, 3, \dots, N$) and the values of RC elements existing in the electrical analogues of a device thermal model of the form of Cauer or Foster network.

The procedure of estimation of the parameters of the device thermal model is automatically realised on the authors' special algorithm implemented in MASTER2. In this algorithm, in the first place the values of the parameters: coefficients a_i , thermal time constant τ_{thi} and the thermal resistance R_{th} are calculated using the measured $Z(t)$ dependence. Next, the values of RC elements in the Foster and Cauer networks are computed.

The considered algorithm is valid in the case when the values of the successive thermal time constants differ from each other at least a few times. Therefore, the dependence $Z(t)$ normalized to R_{th} is a intervally linear function in the lin-log scale. The value of R_{th} results directly from $Z(t)$ at the steady-state. Values of the remaining parameters are calculated by the last-square method used for approximation of any

special function with $Z(t)$ as an argument. At first the longest thermal time constant and the coefficient a_1 corresponding to them are calculated. The values of RC elements of the Foster network are calculated after analytical dependences containing the considered parameters. In turn, to calculate the values of the Cauer network, the operational thermal impedance $Z(s)$ has to be formulated. Next, division of the values of coefficients corresponding to the highest degree of polynomials representing the numerator and the denominator of $Z(s)$ respectively, have to be performed. The value of such division at the k -th iteration is denoted as D^k . In each iteration step, the product of the denominator value and the value of D^k is subtracted from the numerator and then the new numerator and the denominator are transformed to each other. The presented computer tool was verified experimentally by comparing the measured and calculated dependence of $Z(t)$ of the VDMOS transistor IRF840. The very well agreement between measurements and simulations have been obtained.

Keywords: transient thermal impedance, the thermal model, thermal measurements

efficient
calculated
s of the
e values
and the
iteration
of D^k is
to each
calculated
rements

Modelowanie właściwości mikromocowych wzmacniaczy operacyjnych w obszarze nasycenia

JÓZEF STANCLIK

*Politechnika Wrocławska,
Instytut Telekomunikacji, Teleinformatyki i Akustyki,
50-370 Wrocław, ul. Wybrzeże Wyspiańskiego 27,
e-mail: jozef.stanclik@pwr.wroc.pl*

Otrzymano 2005.08.29

Autoryzowano 2005.11.15

W pracy zaproponowano modyfikację makromodeli wzmacniaczy operacyjnych, która prowadzi do uwzględnienia zależności maksymalnego i minimalnego napięcia wyjściowego układu scalonego od jego prądu wyjściowego i temperatury otoczenia. Dzięki wprowadzonej zmianie uzyskano lepszą zgodność maksymalnych napięć wyjściowych makromodelu i układu scalonego, oraz rozszerzenie obszaru warunków pracy (temperatura, rezystancja obciążenia, napięcie wejściowe), w którym ta zwiększona zgodność jest zachowana, co jest bardzo istotne w przypadku mikromocowych wzmacniaczy operacyjnych, szczególnie t.zw. wzmacniaczy „rail-to-rail”, zasilanych bardzo małymi napięciami.

Modyfikacja polega na dodaniu do pierwotnego makromodelu dwóch nieliniowych źródeł napięciowych, sterowanych prądem wyjściowym i prądem proporcjonalnym do temperatury. Przedstawiono sposób wyznaczania współczynników wielomianów, opisujących dodatkowe źródła, oparty o metodę najmniejszych kwadratów. Podstawą do wyznaczania tych współczynników są typowe dane katalogowe układu scalonego. Podano przykłady modyfikacji makromodeli: mikromocowego wzmacniacza operacyjnego TLV2211 i wzmacniacza OPA130, przeznaczonych do stosowania w programie PSpice. Uzyskano zmniejszenie błędów modelowania zależności maksymalnego i minimalnego napięcia wyjściowego od prądu wyjściowego i temperatury otoczenia o około rząd wielkości w stosunku do błędów pierwotnych makromodeli, dostarczanych przez ich wytwórców.

Słowa kluczowe: makromodel, wzmacniacz operacyjny, program PSpice

1. WPROWADZENIE

Symulacja komputerowa układów elektronicznych zawierających wzmacniacze scalone jest zwykle wykonywana po zastąpieniu układów scalonych makromodelami, dzięki czemu ulegają zmniejszeniu czas i koszt obliczeń oraz problemy ze zbieżnością procedur numerycznych stosowanych symulatorów. Do identyfikacji makromodeli używane jest najczęściej specjalistyczne oprogramowanie (np. program PARTS, wchodzący wraz z programem PSpice w skład pakietu Design Center [2]). Uzyskane za ich pomocą makromodele uwzględniają podstawowe parametry i charakterystyki układów scalonych w ustalonej temperaturze.

Mikromocowe wzmacniacze operacyjne mogą pracować przy bardzo małych napięciach zasilających, np. $\pm 1,4$ V (typowo 3 V lub 5 V), a ich maksymalne międzyszczytowe napięcie wyjściowe może mieć wartość niemal równą napięciu zasilającemu (w układach „rail-to-rail” [6]). Maksymalne napięcie wyjściowe makromodeli dostarczanych przez producentów wzmacniaczy nie zależy od prądu wyjściowego układów scalonych. Napięcie to silnie natomiast zależy od temperatury otoczenia. Maksymalne napięcie wyjściowe makromodelu jest zgodne z danymi katalogowymi układu scalonego tylko przy jednej wartości prądu wyjściowego i jednej wartości temperatury.

W przypadku wzmacniaczy operacyjnych ogólnego przeznaczenia, zasilanych napięciami ± 15 V, niezgodność największych, międzyszczytowych wartości napięcia wyjściowego wzmacniacza operacyjnego (typowo 27 V) i jego makromodelu na ogół nie przekracza kilku procent. Maksymalne, międzyszczytowe napięcie wyjściowe mikromocowego wzmacniacza operacyjnego może być znacznie mniejsze. W takim przypadku błąd ułamka wolta może prowadzić do znacznych rozbieżności między wynikami obliczeń symulacyjnych i wynikami pomiarów badanych układów (por. p. 3).

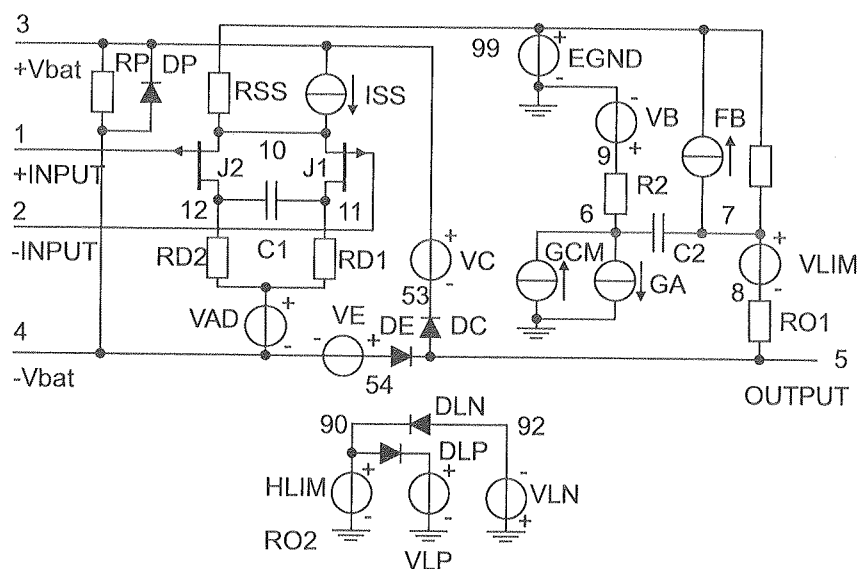
Uwzględnienie zależności napięcia nasycenia od prądu wyjściowego w makromodelach operacyjnych wzmacniaczy mocy jest przedmiotem pracy [3]. Uzależnienie napięcia nasycenia od prądu uzyskano dzięki wprowadzeniu do makromodelu źródeł sterowanych obniżających próg ograniczania napięcia wyjściowego proporcjonalnie do prądu wyjściowego. W przypadku scalonych wzmacniaczy mocy m.cz. podobny efekt dało wprowadzenie rezystancji szeregowych do modeli diod ograniczników napięcia wyjściowego [4]. Oba sposoby są efektywne, gdy zmiany napięcia nasycenia są w przybliżeniu proporcjonalne do zmian prądu wyjściowego. Gdy zależność ta jest nieliniowa można użyć nieliniowych źródeł sterowanych. W pracy [5] aproksymowano nieliniową zależność napięcia nasycenia wzmacniacza operacyjnego od prądu za pomocą wielomianu jednej zmiennej.

W niniejszym opracowaniu uzależniono maksymalne i minimalne napięcie wyjściowe makromodelu od prądu wyjściowego i temperatury otoczenia, którą wcześniej zamieniono na zmienną elektryczną. Do pierwotnego makromodelu dodano dwa nieliniowe źródła napięciowe, sterowane prądem wyjściowym i prądem proporcjonalnym do temperatury. Podano sposób wyznaczania współczynników wielomianów, opisujących dodatkowe źródła, na podstawie typowych danych katalogowych. Zamieszczono przy-

kłady modyfikacji makromodeli mikromocowego wzmacniacza operacyjnego TLV2211 i wzmacniacza ogólnego przeznaczenia OPA130.

2. MODYFIKACJA MAKROMODELU

Strukturę makromodelu mikromocowego wzmacniacza operacyjnego, dostarczanego przez firmę Texas Instruments do stosowania w programie PSpice, pokazano na rys. 1. Jest to schemat makromodelu wzmacniacza wykonanego w technologii CMOS, modelowanego ze standardową dokładnością (LEVEL I). Firma ta opracowała również wersję makromodelu o podwyższonej dokładności (LEVEL II) [6].



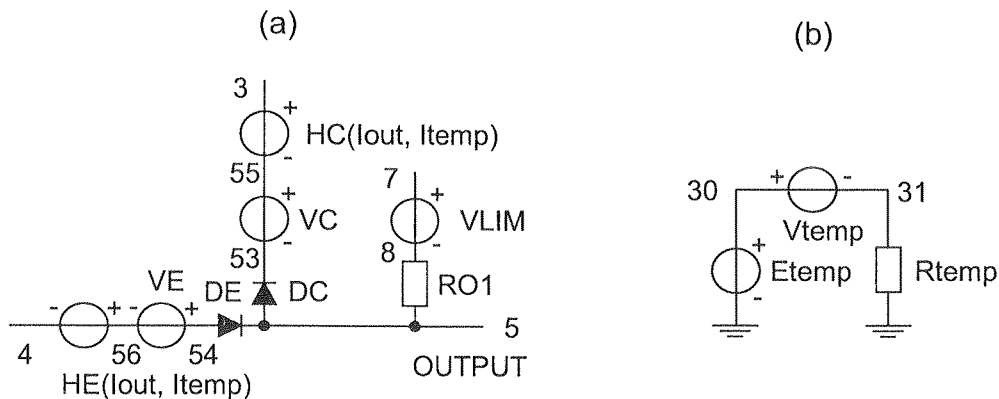
Rys. 1. Struktura makromodelu mikromocowego wzmacniacza operacyjnego

Fig. 1. Structure of macromodel for micropower operational amplifier

Maksymalne i minimalne napięcie wyjściowe w obu wersjach makromodelu jest ograniczane w obwodach ograniczników napięcia, złożonych z idealnych diod półprzewodnikowych DC i DE , oraz źródeł pomocniczych VC i VE . Napięcia progowe są ustalane przez napięcia zasilające $+V_{bat}$ i $-V_{bat}$ oraz siły elektromotoryczne źródeł VC i VE . Maksymalne i minimalne napięcie wyjściowe makromodelu praktycznie nie zależy od prądu wyjściowego. Natomiast zmiana temperatury wpływa silnie na te napięcia w wyniku zmiany spadków napięć na diodach DC i DE (o około $2,7 \text{ mV/}^{\circ}\text{C}$). Pociąga to za sobą zmiany napięć progowych ograniczników napięcia wyjściowego przewyższające $0,4 \text{ V}$ w pełnym zakresie zmian temperatury od -40°C do $+125^{\circ}\text{C}$.

Napięcie wyjściowe rzeczywistego mikromocowego wzmacniacza operacyjnego zmienia się w zależności od prądu wyjściowego i zależy od temperatury otoczenia. W danych katalogowych jest podawana zależność maksymalnego (V_{OH}) i minimalnego (V_{OL}) napięcia wyjściowego od prądu wyjściowego (odpowiednio I_{OH} i I_{OL}) przy kilku wartościach temperatury otoczenia (por. dane katalogowe wzmacniacza operacyjnego TLV2211 [6]).

Uzależnienie maksymalnej i minimalnej wartości napięcia wyjściowego makromodelu od prądu wyjściowego i temperatury można uzyskać w wyniku dodania do obwodów ograniczników napięcia wyjściowego makromodelu nieliniowych źródeł napięciowych HC i HE , sterowanych prądem wyjściowym i temperaturą, przesuwających odpowiednio napięcia progowe ograniczników (rys. 2a). Do próbkowania prądu wyjściowego makromodelu można użyć źródła V_{lim} . Prąd tego źródła jest praktycznie równy prądowi wyjściowemu, $I_{out} = I(V_{lim})$, bo prądy diod DC i DE są o kilka rzędów mniejsze.



Rys. 2. Schematy: (a) ograniczników napięcia w zmodyfikowanym makromodelu, (b) obwodu przekształcającego temperaturę w zmienną elektryczną

Fig. 2. Diagrams: (a) voltage limiters in modified macromodel, (b) circuit converting temperature to electrical quantity

W programie PSpice temperatura $TEMP$ jest parametrem zadawanym przez użytkownika (standardowo 27°C). Zamiana tego parametru na zmienną elektryczną jest możliwa dzięki funkcji $VALUE\{x\}$. Jest ona realizowana w obwodzie pokazanym na rys. 2b, złożonym ze sterowanego źródła napięciowego $Etemp$, opisanego wzorem (1), źródła autonomicznego $Vtemp$ o wartości temperatury standardowej $Vtemp = 27$ i rezystora $Rtemp$, którego rezystancję przyjęto równą $Rtemp = 1\text{k}\Omega$.

$$Etemp = VALUE\{TEMP\} \quad (1)$$

Fragment listy połączeń zmodyfikowanego makromodelu, w którym różnica temperatury otoczenia i temperatury standardowej 27°C jest zamieniana na prąd $I_{temp} = I(Vtemp)$, używany do sterowania źródłami HC i HE , zamieszczono w tabeli 1.

Tabela 1

Fragment listy połączeń zmodyfikowanego makromodelu

Part of modified macromodel netlist

ETEMP	30	0	VALUE {TEMP}
VTEMP	30	31	27
RTEMP	31	0	1K

Źródła HC i HE aproksymujące nieliniowe zależności sił elektromotorycznych od prądów I_{out} i I_{temp} są opisane zależnościami (2) i (3).

$$HC = a_{0C} + a_{1C}I_{out} + a_{2C}I_{temp} + a_{3C}I_{out}^2 + a_{4C}I_{out}I_{temp} + a_{5C}I_{temp}^2 + \dots \quad (2)$$

$$HE = a_{0E} + a_{1E}I_{out} + a_{2E}I_{temp} + a_{3E}I_{out}^2 + a_{4E}I_{out}I_{temp} + a_{5E}I_{temp}^2 + \dots \quad (3)$$

Stopień wielomianu może być dowolny, jednak stosowanie wielomianów stopnia wyższego niż drugi nie jest celowe, bo choć zwiększa się dokładność aproksymacji, to rośnie liczba parametrów modelu.

Podstawą do wyznaczenia parametrów źródła HC są wartości różnic maksymalnych napięć wyjściowych: wzmacniacza operacyjnego V_{OHIC} (odczytanych z danych katalogowych) i pierwotnego makromodelu V_{OHmm} (z symulacji komputerowej), wyznaczone dla n wartości prądu wyjściowego I_{OH} i m wartości temperatur otoczenia T_a :

$$\Delta V_{OHij} = V_{OHmmij} - V_{OHICij} = f(I_{OHi}, T_{aj} - 27), \quad 1 \leq i \leq n, 1 \leq j \leq m \quad (4)$$

(przyjęto standardową temperaturę otoczenia 27°C). Analogicznie podstawą do wyznaczenia parametrów źródła HE jest zależność (5).

$$\Delta V_{OLkj} = V_{OLmmkj} - V_{OLICKj} = f(I_{OLk}, T_{aj} - 27), \quad 1 \leq i \leq n, 1 \leq j \leq m \quad (5)$$

Współczynniki wielomianów HC i HE drugiego stopnia są wyznaczane w wyniku rozwiązania układów równań, które minimalizują błędy ε_{ij} oraz ε_{kj} , dane zależnościami (6) i (7), w sensie przyjętego kryterium (zwykle średniokwadratowego lub minimaxowego) ze względu na parametry $a_{0C}, \dots, a_{5C}$ i $a_{0E}, \dots, a_{5E}$ (odpowiednio).

$$\begin{aligned} \varepsilon_{ij}(a_{0C}, a_{1C}, \dots, a_{5C}) = & \Delta V_{OLkj} - \{a_{0C} + a_{1C}I_{OLk} + a_{2C}I(Vtemp_j) + \\ & + a_{3C}[I_{OLk}]^2 + a_{4C}I_{OLk}I(Vtemp_j) + a_{5C}[I(Vtemp_j)]^2\}, \end{aligned} \quad (6)$$

$$1 \leq i \leq n, 1 \leq j \leq m$$

$$\begin{aligned} \varepsilon_{kj}(a_{0E}, a_{1E}, \dots, a_{5E}) = \Delta V_{OHkj} - \{a_{0E} + a_{1E}I_{OHk} + a_{2E}I(Vtemp_j) + \\ + a_{3E}[I_{OHk}]^2 + a_{4E}I_{OHk}I(Vtemp_j) + a_{5E}[I(Vtemp_j)]^2\}, \end{aligned} \quad (7)$$

$$1 \leq k \leq n, 1 \leq j \leq m$$

W Dodatku przedstawiono przebieg odpowiednich obliczeń dla średniokwadratowego kryterium błędu. Zastosowanie jednego z uniwersalnych programów matematycznych (Matlab, Mathcad, Mathematica, itp.) ułatwia wyznaczenie współczynników wielomianów (2) i (3).

3. PRZYKŁAD

Zaproponowaną modyfikację wprowadzono do makromodelu mikromocowego wzmacniacza operacyjnego TLV2211. Jego listę połączeń w formacie stosowanym w bibliotekach programu PSpice zamieszczono w tabeli 2 [6] (w pierwotnym makromodelu nie występują elementy zapisane tłustym drukiem). Zależność maksymalnego napięcia wyjściowego V_{OHIC} od prądu wyjściowego I_{OH} przy kilku temperaturach T_a jest podawana w danych katalogowych układu scalonego w postaci wykresów pokazanych na rys. 3.

Na rysunku tym wkreślono również linią przerywaną zależności napięcia V_{OHmm} od prądu wyznaczone na drodze symulacji komputerowej przy użyciu pierwotnego makromodelu. Na podstawie wykresów odczytano wartości różnic ΔV_{OH} dla przyjętych prądów i temperatur, a następnie wyznaczono współczynniki wielomianu drugiego stopnia (2). Uzyskano wielomian o postaci:

$$HC = -0,15 + 750 \cdot I(Vlim) - 2,74 \cdot I(Vtemp) - 4,2 \cdot 10^3 \cdot I(Vlim) \cdot I(Vtemp) \quad (8)$$

W podobny sposób wyznaczono współczynniki wielomianu (3):

$$HE = -0,15 - 290 \cdot I(Vlim) - 2,69 \cdot I(Vtemp) - 1,7 \cdot 10^3 \cdot I(Vlim) \cdot I(Vtemp) \quad (9)$$

Zmiany wprowadzone do makromodelu wzmacniacza TLV2211 zaznaczono tłustym drukiem w jego zapisie, zamieszczonym w tabeli 2. Przy użyciu zmienionego makromodelu wykonano obliczenia maksymalnego napięcia wyjściowego V_{OHmod} i wkreślono je linią kropkowaną na rys. 3. Widać, że o ile maksymalne napięcie wyjściowe, obliczone z użyciem pierwotnego makromodelu, różni się od danych katalogowych układu scalonego nawet o ponad 400 mV, to po modyfikacji maksymalny błąd nie przekracza 10 mV. Również w przypadku minimalnego napięcia wyjściowego V_{OL} uzyskano zmniejszenie największego błędu o więcej niż rząd wielkości.

Podobnej modyfikacji poddano makromodel wzmacniacza operacyjnego OPA130 [1], pochodzący z biblioteki modeli programu PSpice [2]. Zastosowano źródła sterowane opisane wielomianami (10) i (11).

$$\begin{aligned} HC = -0,1 + 34,5 \cdot I(Vlim) - 2,36 \cdot I(Vtemp) + 12,7 \cdot 10^3 \cdot [I(Vlim)]^2 + \\ + 741 \cdot I(Vlim) \cdot I(Vtemp) \end{aligned} \quad (10)$$

Tabela 2

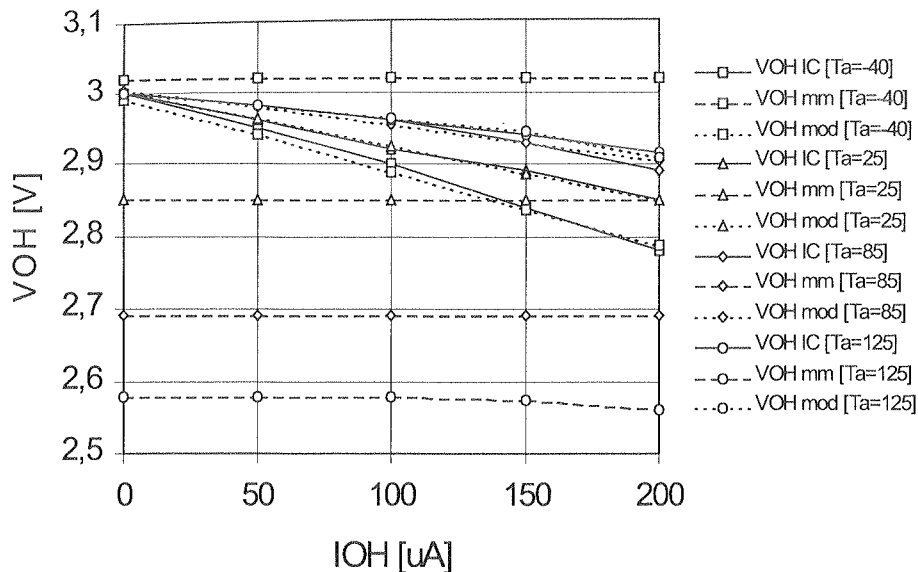
Zmodyfikowany makromodel mikromocowego wzmacniacza operacyjnego TLV2211

Modified macromodel of TLV2211 micropower operational amplifier

```

* TLV2211 micropower operational amplifier macromodel
+ subcircuit
* connections:      non-inverting input
*                  | inverting input
*                  | | positive power supply
*                  | | | negative power supply / ground
*                  | | | | output
*                  | | | | |
.SUBCKT TLV2211 1 2 3 4 5
J1 11 2 10 JX
J2 12 1 10 JX
RD1 60 11 106.1E+03
RD2 60 12 106.1E+03
C1 11 12 8.86E-12
C2 6 7 50.0E-12
DC 5 53 DX
DE 54 5 DX
DLP 90 91 DX
DLN 92 90 DX
DP 4 3 DX
EGND 99 0 POLY(2) 3 0 4 0 0 .5 .5
FB 7 99 POLY(5) VB VC VE VLP VLN 0 4.29E6 -6E6
+6E6 6E6 -6E6
ISS 3 10 DC 1.250E-06
RSS 10 99 160.0E+06
GA 6 0 11 12 9.425E-06
R2 6 9 1.000E+05
GCM 0 6 10 99 1320.2E-12
RO2 7 99 150
RP 3 4 419.2E3
VAD 60 4 -.5
VB 9 0 DC 0
RO1 8 5 50
VC 55 53 DC .55 ;VC 3 53 DC .55
ETEMP 30 0 VALUE {TEMP}
VTEMP 30 31 27
RTEMP 31 0 1K
HC 3 55 POLY(2) VLIM VTEMP -0.15 750 -2.74 0 -4.2E3
HE 56 4 POLY(2) VLIM VTEMP -0.15 -290 -2.69 0 -1.7E3
VE 54 56 DC .55 ;VE 54 4 DC .55
VLIM 7 8 DC 0
VLP 91 0 DC .1
VLN 0 92 DC 2.6
HLIM 90 0 VLIM 1.000E+03
.MODEL DX D (IS=8.000E-16)
.MODEL JX PJF (IS=500.0E-15 BETA=166E-6 VTO=-.004)
.ENDS

```



Rys. 3. Zależność maksymalnego napięcia wyjściowego wzmacniacza operacyjnego TLV2211 (VOH IC) i jego makromodeli: pierwotnego (VOHmm) i zmodyfikowanego (VOHmod) od prądu wyjściowego (IOH) i temperatury (Ta)

Fig. 3. Maximum output voltage versus output current (IOH) and ambient temperature (Ta) for TLV2211 operational amplifier (VOH IC) and its macromodels: original (VOHmm) and modified (VOHmod)

$$HE = 19,7 \cdot I(Vlim) - 2,36 \cdot I(Vtemp) + 22,9 \cdot 10^3 \cdot [I(Vlim)]^2 + 2,28 \cdot 10^3 \cdot I(Vlim) \cdot I(Vtemp) \quad (11)$$

Zmiany wprowadzone do makromodelu wzmacniacza OPA130 zaznaczono tłustym drukiem w jego zapisie, zamieszczonym w tabeli 3. Modyfikacja doprowadziła do zmniejszenia błędów o około jeden rząd wielkości w zakresie zmian prądu wyjściowego od 0 do 10 mA i w przedziale temperatur od 25°C do 125°C.

4. PODSUMOWANIE I WNIOSKI KOŃCOWE

Zaproponowana modyfikacja makromodeli wzmacniaczy operacyjnych gwarantuje uwzględnienie zależności maksymalnego i minimalnego napięcia wyjściowego od prądu wyjściowego układu scalonego i temperatury otoczenia. Makromodele dostarczane przez wytwórców układów scalonych i dostępne w bibliotekach programu PSpice nie mają takiej właściwości. Wysoka zgodność maksymalnych napięć wyjściowych układu scalonego i jego makromodelu jest zachowana w szerszym zakresie zmian tempera-

Tabela 3

Zmodyfikowany makromodel wzmacniacza operacyjnego OPA130

Modified macromodel of OPA130 operational amplifier

```

* OPA130 OPERATIONAL AMPLIFIER "MACROMODEL" SUBCIRCUIT
*
.SUBCKT OPA130
+ 1 ; non-inverting input
+ 2 ; inverting input
+ 3 ; positive power supply
+ 4 ; negative power supply
+ 5 ; output
*
C1 11 12 4.330E-12
C2 6 7 15.00E-12
CSS 10 99 1.000E-30
DC 5 53 DX
DE 54 5 DX
DLP 90 91 DX
DLN 92 90 DX
DP 4 3 DX
EGND 99 0 POLY(2) (3,0) (4,0) 0 .5 .5
FB 7 99 POLY(5) VB VC VE VLP VLN 0 2.386E9 -2E9
+2E9 2E9 -2E9
GA 6 0 11 12 94.25E-6
GCM 0 6 10 99 530.1E-12
ISS 3 10 DC 30.00E-6
HLIM 90 0 VLIM 1E3
J1 11 2 10 JX
J2 12 1 10 JX
R2 6 9 100.0E3
RD1 4 11 10.61E3
RD2 4 12 10.61E3
RO1 8 5 50
RO2 7 99 25
RP 3 4 56.25E3
RSS 10 99 6.667E6
VB 9 0 DC 0
VC 55 53 DC 1.500 ;VC 3 53 DC 1.500
HC 3 55 POLY(2) VLIM VTEMP -.1 34.5 -2.36 12.7E3 741
HE 56 4 POLY(2) VLIM VTEMP 0 19.7 -2.36 22.9E3 -2.28E3
ETEMP 30 0 VALUE {TEMP}
VTEMP 30 31 25
RTEMP 31 0 1K
VE 54 56 DC 1 ;VE 54 4 DC 1
VLIM 7 8 DC 0
VLP 91 0 DC 18
VLN 0 92 DC 18
.MODEL DX D(IS=800.0E-18)
.MODEL JX PJF(IS=2.500E-12 BETA=296.1E-6 VTO=-1)
.ENDS

```

tury, rezystancji obciążenia, oraz napięć wejściowych, niż w przypadku makromodeli oryginalnych.

Stopień złożoności makromodelu zwiększa się nieznacznie, bowiem modyfikacja polega na uzupełnieniu ograniczników napięcia wyjściowego o dwa nieliniowe źródła napięciowe, sterowane prądem wyjściowym makromodelu i pośrednio temperaturą otoczenia. Ze względu na liczbę dodatkowych parametrów makromodelu do opisu źródeł są stosowane wielomiany stopnia co najwyżej drugiego. Wprowadzenie modyfikacji do listy połączeń makromodelu polega na zmianie dwóch i dopisaniu pięciu wierszy do zapisu pierwotnego makromodelu w bibliotece symulatora układów elektronicznych. Typowe dane katalogowe układu scalonego są wystarczającą podstawą do obliczenia współczynników wielomianów (nie są konieczne dodatkowe pomiary układu scalonego).

Zamieszczone przykłady modyfikacji makromodeli mikromocowego wzmacniacza operacyjnego TLV2211 i wzmacniacza OPA130, świadczą o przydatności zmodyfikowanych makromodeli do zastosowań, w których układ scalony może pracować w szerokim zakresie zmian temperatury, rezystancji obciążenia i wysterowania, a wymagana jest duża zgodność maksymalnego i minimalnego napięcia wyjściowego układu scalonego i jego makromodelu. W wyniku modyfikacji, największy błąd modelowania zależności maksymalnego i minimalnego napięcia wyjściowego układu scalonego od prądu wyjściowego i temperatury uległ zmniejszeniu o około rząd wielkości w stosunku do błędu uzyskiwanego w przypadku stosowania pierwotnego makromodelu.

5. BIBLIOGRAFIA

1. BURR-BROWN IC Data Book — Linear Products 1996. Burr-Brown Corporation, 1996/1997, pp. 2.53-2.59.
2. MicroSim PSpice A/D. Ver.6.3. Reference Manual. MicroSim Corporation, April 1996.
3. J. Stanclik: *Modelowanie napięcia nasycenia wzmacniaczy operacyjnych*. Materiały VIII KK KOWBAN'2001, Wrocław — Świeradów Zdrój, 25-27.10.2001, ss. 163-168.
4. J. Stanclik: *Zwiększenie dokładności makromodelu skalonych wzmacniaczy mocy małej częstotliwości*. Kwartalnik Elektroniki i Telekomunikacji, 2001, T. 47, z. 4, ss. 463-476.
5. J. Stanclik: *Modelowanie zależności maksymalnego napięcia wyjściowego wzmacniaczy operacyjnych od prądu wyjściowego*. Kwartalnik Elektroniki i Telekomunikacji, 2004, T. 50, z. 3, ss. 411-427.
6. Texas Instruments, Designer's Guide and Databook, Mixed-Signal & Analog, Sept. 1996, (CD-ROM).

6. DODATEK

Wyznaczenie wartości parametrów $a_0, \dots, a_5$ funkcji (2) lub (3), minimalizujących błędy aproksymacji, odpowiednio (6) lub (7), w sensie kryterium średniokwadratowego przebiega w sposób przedstawiony poniżej (ze względów praktycznych stopień wielomianu ograniczono do drugiego).

Dane są wartości dwóch zmiennych niezależnych x_i i y_j oraz zmiennej zależnej $z_{ij} = z(x_i, y_j)$ w n, m oraz nm (odpowiednio) punktach:

$$x_i, y_j, z_{ij}, \quad i \in \{1, 2, \dots, n\}, j \in \{1, 2, \dots, m\} \quad (D1)$$

Zależność $z(x, y)$ przybliżamy wielomianem kwadratowym dwóch zmiennych

$$f(x, y) = a_0 + a_1x + a_2y + a_3x^2 + a_4xy + a_5y^2 \quad (D2)$$

w taki sposób, aby zminimalizować średniokwadratowy błąd aproksymacji

$$S(a_0, a_1, \dots, a_5) = \sum_{i=1}^n \sum_{j=1}^m \left(z_{ij} - a_0 - a_1x_i - a_2y_j - a_3x_i^2 - a_4x_iy_j - a_5y_j^2 \right)^2 \quad (D3)$$

W minimum pochodne cząstkowe się zerują:

$$\begin{aligned} \frac{\delta S}{\delta a_0} &= \sum_{i=1}^n \sum_{j=1}^m 2 \left(z_{ij} - a_0 - a_1x_i - a_2y_j - a_3x_i^2 - a_4x_iy_j - a_5y_j^2 \right) (-1) = 0 \\ \frac{\delta S}{\delta a_1} &= \sum_{i=1}^n \sum_{j=1}^m 2 \left(z_{ij} - a_0 - a_1x_i - a_2y_j - a_3x_i^2 - a_4x_iy_j - a_5y_j^2 \right) (-x_i) = 0 \\ \frac{\delta S}{\delta a_2} &= \sum_{i=1}^n \sum_{j=1}^m 2 \left(z_{ij} - a_0 - a_1x_i - a_2y_j - a_3x_i^2 - a_4x_iy_j - a_5y_j^2 \right) (-y_j) = 0 \\ \frac{\delta S}{\delta a_3} &= \sum_{i=1}^n \sum_{j=1}^m 2 \left(z_{ij} - a_0 - a_1x_i - a_2y_j - a_3x_i^2 - a_4x_iy_j - a_5y_j^2 \right) (-x_i^2) = 0 \\ \frac{\delta S}{\delta a_4} &= \sum_{i=1}^n \sum_{j=1}^m 2 \left(z_{ij} - a_0 - a_1x_i - a_2y_j - a_3x_i^2 - a_4x_iy_j - a_5y_j^2 \right) (-x_iy_j) = 0 \\ \frac{\delta S}{\delta a_5} &= \sum_{i=1}^n \sum_{j=1}^m 2 \left(z_{ij} - a_0 - a_1x_i - a_2y_j - a_3x_i^2 - a_4x_iy_j - a_5y_j^2 \right) (-y_j^2) = 0 \end{aligned} \quad (D4)$$

Mnożąc układ (D4) przez $\frac{1}{2}$ i rozdzielając sumy składników otrzymujemy układ równań normalnych (D5), który można rozwiązać za pomocą dowolnego programu rozwiązującego układy równań liniowych.

$$x = A b \quad (D5)$$

gdzie:

$$x^T = [a_0, a_1, a_2, a_3, a_4, a_5]$$

$$b^T = \left[\sum_{i=1}^n \sum_{j=1}^m z_{ij}, \sum_{i=1}^n \sum_{j=1}^m x_i z_{ij}, \sum_{i=1}^n \sum_{j=1}^m y_j z_{ij}, \sum_{i=1}^n \sum_{j=1}^m x_i^2 z_{ij}, \sum_{i=1}^n \sum_{j=1}^m x_i y_j z_{ij}, \sum_{i=1}^n \sum_{j=1}^m y_j^2 z_{ij} \right]$$

$$A = \begin{bmatrix} nm & m \sum_{i=1}^n x_i & n \sum_{j=1}^m y_j & m \sum_{i=1}^n x_i^2 & \sum_{i=1}^n x_i \sum_{j=1}^m y_j & n \sum_{j=1}^m y_j^2 \\ m \sum_{i=1}^n x_i & m \sum_{i=1}^n x_i^2 & \sum_{i=1}^n x_i \sum_{j=1}^m y_j & m \sum_{i=1}^n x_i^3 & \sum_{i=1}^n x_i^2 \sum_{j=1}^m y_j & \sum_{i=1}^n x_i \sum_{j=1}^m y_j^2 \\ n \sum_{j=1}^m y_j & \sum_{i=1}^n x_i \sum_{j=1}^m y_j & n \sum_{j=1}^m y_j^2 & \sum_{i=1}^n x_i^2 \sum_{j=1}^m y_j & \sum_{i=1}^n x_i \sum_{j=1}^m y_j^2 & n \sum_{j=1}^m y_j^3 \\ m \sum_{i=1}^n x_i^2 & m \sum_{i=1}^n x_i^3 & \sum_{i=1}^n x_i^2 \sum_{j=1}^m y_j & m \sum_{i=1}^n x_i^4 & \sum_{i=1}^n x_i^3 \sum_{j=1}^m y_j & \sum_{i=1}^n x_i^2 \sum_{j=1}^m y_j^2 \\ \sum_{i=1}^n x_i \sum_{j=1}^m y_j & \sum_{i=1}^n x_i^2 \sum_{j=1}^m y_j & \sum_{i=1}^n x_i \sum_{j=1}^m y_j^2 & \sum_{i=1}^n x_i^3 \sum_{j=1}^m y_j & \sum_{i=1}^n x_i^2 \sum_{j=1}^m y_j^2 & \sum_{i=1}^n x_i \sum_{j=1}^m y_j^3 \\ n \sum_{j=1}^m y_j^2 & \sum_{i=1}^n x_i \sum_{j=1}^m y_j^2 & n \sum_{j=1}^m y_j^3 & \sum_{i=1}^n x_i^2 \sum_{j=1}^m y_j^2 & \sum_{i=1}^n x_i \sum_{j=1}^m y_j^3 & n \sum_{j=1}^m y_j^4 \end{bmatrix}$$

J. STANCLIK

MODELLING OF MICROPOWER OPERATIONAL AMPLIFIERS PROPRIETIES IN SATURATION REGION

S u m m a r y

The paper proposes specific modification to macromodels of operational amplifiers causing that these models reflect additionally how IC's maximum and minimum output voltages depend on its output current and on ambient temperature. This modification caused better consistency between maximum output voltages of the macromodel and those of the IC; extended was also the area of operating conditions (temperature, load resistance, input voltage) where such improved consistency is maintained. This is of special importance in case of micropower operational amplifiers, especially for so called rail-to-rail amplifiers operated at very low supply voltages.

The modification consists in supplementing the primary macromodel with two non-linear voltage sources controlled by output current and by a current proportional to temperature. The least squares method is presented to determine the coefficients of polynomials representing the additional sources. Typical IC's data-sheet information is used to find out these coefficients. Exemplary modifications are shown for macromodels of TLV2211 micropower operational amplifier, and OPA130 amplifier, applicable for use in PSpice Programme. The relation of maximum and/or minimum output voltage versus output current and ambient temperature is now modelled with accuracy several times better than that attainable for primary macromodel.

Keywords: macromodel, operational amplifier, PSpice programme

Ocena przydatności przetworników pomiarowych wielkości elektrycznych do pomiaru sygnałów odkształconych

KRZYSZTOF PACHOLSKI

*Instytut Elektrotechniki Teoretycznej, Metrologii i Materiałoznawstwa,
Politechnika Łódzka,
ul. Stefanowskiego 18/22, 90-924 Łódź,
e-mail: kpacholski@go2.pl*

Otrzymano 2005.10.14

Autoryzowano 2005.12.21

W artykule zaprezentowano sposób wyznaczania parametrów szybkościowych przetworników pomiarowych mocy oraz przetworników napięciowych wartości średniej i skutecznej mierzących szybkozmienne sygnały odkształcone. Parametrami tymi są: dopuszczalna szybkość zmian sygnału wejściowego S_d oraz częstotliwość graniczna f_p pasma przetwarzania mocy wejściowego sygnału sinusoidalnego. Parametrów S_d i f_p nie uwzględnia, zalecana w obowiązującej normie, metoda oceny właściwości metrologicznych wymienionych przetworników.

Słowa kluczowe: przetworniki napięciowe, przetworniki pomiarowe mocy, parametry szybkościowe, sygnały odkształcone

1. WPROWADZENIE

Sygnały pomiarowe występujące w obwodach zasilanych za pośrednictwem przekształtników energoelektronicznych są coraz częściej sygnałami odkształconymi, o stromych zboczach i częstotliwości harmonicznej podstawowej należącej do pasma częstotliwości akustycznych [5]. Dotychczas przy doborze elektronicznych przetworników pomiarowych mierzących wielkości elektryczne w tych obwodach nie uwzględniano wpływu przesterowania szybkościowego na ich właściwości metrologiczne. Naturę fizyczną przesterowania szybkościowego oraz jego wpływ na właściwości metrologiczne elektronicznych przetworników napięciowych wartości średniej i skutecznej oraz przetworników mocy autor opisał w pracach [6, 7, 8, 9, 10, 11]. W pracach [11, 13] uzasadnił również celowość rozszerzenia zbioru parametrów charakteryzujących przy-

datność przetworników wielkości elektrycznych do pomiaru sygnałów odkształconych o dwa nowe składniki, jakimi są parametry szybkościowe: S_d i f_p . Znajomość tych parametrów pozwoli uniknąć przesterowania szybkościowego przetworników, które jest niejednokrotnie przyczyną błędu pomiaru o dużej wartości [6, 7, 8, 9, 10, 11, 13].

Parametr S_d wykorzystać można przy doborze przetwornika, odpowiednio do szybkości zmian S_x mierzonego sygnału ($S_d > S_x$). Znajomość częstotliwości f_p niezbędna jest do wyznaczenia częstotliwości f_{pD} ($f_{pD} \leq f_p$) ograniczającej pasmo przetwarzania, w którym sygnał odkształcony nie powoduje przesterowania szybkościowego przetwornika [11].

W artykule przedstawiono sposób wyznaczania parametrów szybkościowych (S_d i f_p) przetworników pomiarowych mocy oraz przetworników napięciowych wartości średniej i skutecznej. Wykorzystano do tego celu odpowiednio ukształtowany sygnał testujący charakteryzujący się maksymalną szybkość zmian równą dopuszczalnej szybkości zmian S_d sygnału wejściowych badanego przyrządu, gdy częstotliwość tego sygnału ma wartość graniczną f_{pD} [11]. Przy odkształconym sygnale wejściowym o częstotliwości równej f_{pD} składowa błędu przetwornika, której przyczyną jest przesterowanie szybkościowe, ma wartość równą zero. Błąd ten autor nazwał błędem przesterowania i zdefiniował zależnością [11, 13]

$$\delta_p^{df} = \frac{\Delta_p}{Y} = \frac{Y_p}{Y} - 1 \quad (1)$$

W zal. (1) Y_p i Y oznaczają odpowiednio sygnał wyjściowy przetwornika przesterowanego szybkościowo oraz sygnał wyjściowy tego przetwornika, gdy pracuje on bez przesterowania szybkościowego. Wartości Y_p i Y sygnału wyjściowego wyznaczone muszą być przy zachowaniu stałej wartości wielkości mierzonej przez przetwornik.

Przed wyznaczeniem zerowej wartości błędu δ_p należy zmienić parametry odkształconego sygnału testującego tak, aby sygnał ten był przetwarzany przez badany przetwornik z błędem przesterowania o wartości różnej od zera.

Z definicji błędu przesterowania (zal. 1) wynika, że do wyznaczenia tego błędu nie należy stosować przyrządu wzorcowego. Powodem tego jest niemożliwość doboru przyrządu wzorcowego charakteryzującego się małosygnałowymi parametrami dynamicznymi, które są równe małosygnałowym parametrom dynamicznym przetwornika badanego. Ponadto dopuszczalna szybkość zmian S_d sygnału wejściowego tego przetwornika musi być bliska nieskończoności. Wymienione przyczyny wskazują na konieczność wykorzystania, do wyznaczenia błędu przesterowania (zal. 1), różnicy wskazań badanego przyrządu. W tym celu należy wyznaczyć wartość wskazania Y_p , dla sygnału testującego charakteryzującego się szybkością zmian S_{xm} większą od wartości parametru S_d przetwornika badanego, oraz wartość wskazania Y dla sygnału testującego przetwarzanego bez przesterowania szybkościowego, którego maksymalna szybkość zmian jest mniejsza od S_d . Synteza kształtu takich sygnałów testujących jest przedmiotem dalszych rozważań.

2. WYZNACZANIE PARAMETRÓW SZYBKOSCIOWYCH PRZETWORNIKÓW NAPIĘCIOWYCH

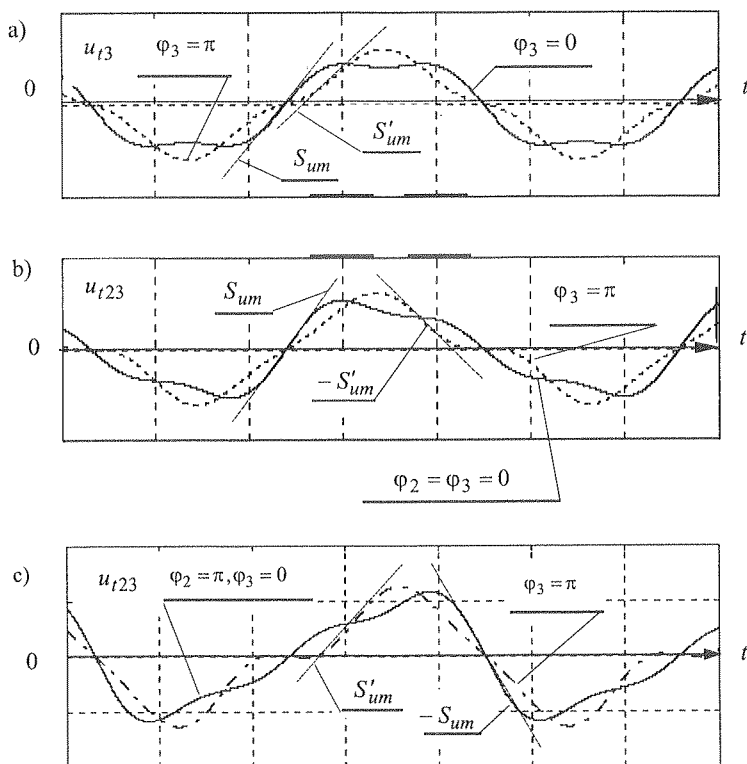
2.1. SYNTEZA SYGNAŁU TESTUJĄCEGO

Dokładność wyznaczania parametrów szybkościowych zależy jest od parametrów stosowanych w badaniach odkształconych sygnałów testujących. Aby dokładność wyznaczania tych parametrów nie była zależna od małosygnałowych błędów częstotliwościowych przetwornika badanego, sygnały testujące muszą charakteryzować się taką samą częstotliwością oraz zgodnością przebiegu obwiedni widma amplitudowego. W celu zminimalizowania wpływu ograniczonej szerokości pasma przetwarzania na błąd wyznaczania parametru S_d badanego przetwornika użyteczny zakres widma sygnałów testujących [11, 12] należeć musi do pasma przetwarzania tego przetwornika. Do grupy sygnałów spełniających ww. wymagania należy sygnał u_{t3} odkształcony trzecią harmoniczną (rys. 1a) i sygnały u_{t23} stanowiące sumę harmonicznnej podstawowej oraz drugiej i trzeciej harmonicznnej (rys. 1b i c) [11, 13].

Niezbędną do występowania błędu przesterowania zmianę wartości maksymalnej szybkości tych sygnałów uzyskać można zmieniając fazę trzeciej harmonicznnej z 0 rad na π rad (rys. 1), bez zmiany ich częstotliwości [11, 12].

Sygnały u_{t3} oraz u_{t23} można wykorzystać do wyznaczania parametrów szybkościowych przetworników, gdy udział h_k ($k = 2, 3$) harmonicznnych należy do przedziału wartości, w którym pomiędzy błędem przesterowania i parametrem h_k zachodzi jednoznaczny związek [11]. Przy spełnieniu tego warunku występuje jednoznaczna zależność błędu przesterowania od maksymalnej szybkości zmian S_{um} sygnałów u_{t3} i u_{t23} , z trzecią harmoniczną o zerowej fazie ($\varphi_3 = 0$ rad). Do wyznaczenia górnej granicznej wartości udziału h_k ($k = 2, 3$) harmonicznnych składowych, sygnału u_{t3} oraz sygnałów u_{t23} (rys. 1), autor wykorzystał wyznaczone numerycznie charakterystyki zależności błędu przesterowania δ_p (zał. 1) przetwornik napięciowych wartości średniej i skutecznej od wartości parametru h_k odpowiedniej harmonicznnej [11]. Przy wyznaczaniu charakterystyk $\delta_p = (h_k)$ przyczyną błędu przesterowania była zmiana fazy, z 0 na π rad trzeciej harmonicznnej sygnałów u_{t3} i u_{t23} o stałej wartości względnej częstotliwości $\frac{f_1}{f_{pD}} \cdot (f_1, f_{pD})$ — (f_1, f_{pD} — oznaczają odpowiednio częstotliwość harmonicznnej podstawowej sygnałów u_{t3} i u_{t23} oraz graniczną częstotliwość pasma przetwarzania mocy tych sygnałów testujących przez badany przetwornik). Charakterystyki $\delta_p = (h_k)$ wyznaczono zakładając, że przyczyną przesterowania przetworników napięciowych jest ograniczona szybkość zmian S sygnału wyjściowego wzmacniaczy operacyjnych zastosowanych w wejściowym bloku przemiennoprądowym tych przetworników [11].

Przebieg charakterystyk $\delta_p = (h_3)$, wyznaczonych dla przetwornika wartości średniej oraz dla przetwornika wartości skutecznej, gdy na ich wejściu występował sygnał u_{t3} (rys. 1a), przedstawiają odpowiednio rys. 2a i b [11].



Rys. 1. Sygnał u_{t3} odkształcony trzecią harmoniczną (a) oraz sygnały u_{t23} odkształcone drugą i trzecią harmoniczną (b) i (c) (S_{um} , S'_{um} oznaczają maksymalne szybkości zmian sygnałów ($S_{um} > S'_{um}$); h_2 , h_3 — oznaczają udział drugiej i trzeciej harmonic-znej)

Fig. 1. Signal u_{t3} deformed with third harmonic (a) and signals u_{t23} deformed with harmonics: second and third (b) i (c) (S_{um} , S'_{um}) signify maximum rate of signals ($S_{um} > S'_{um}$); h_2 , h_3 signify share of the second and third harmonic)

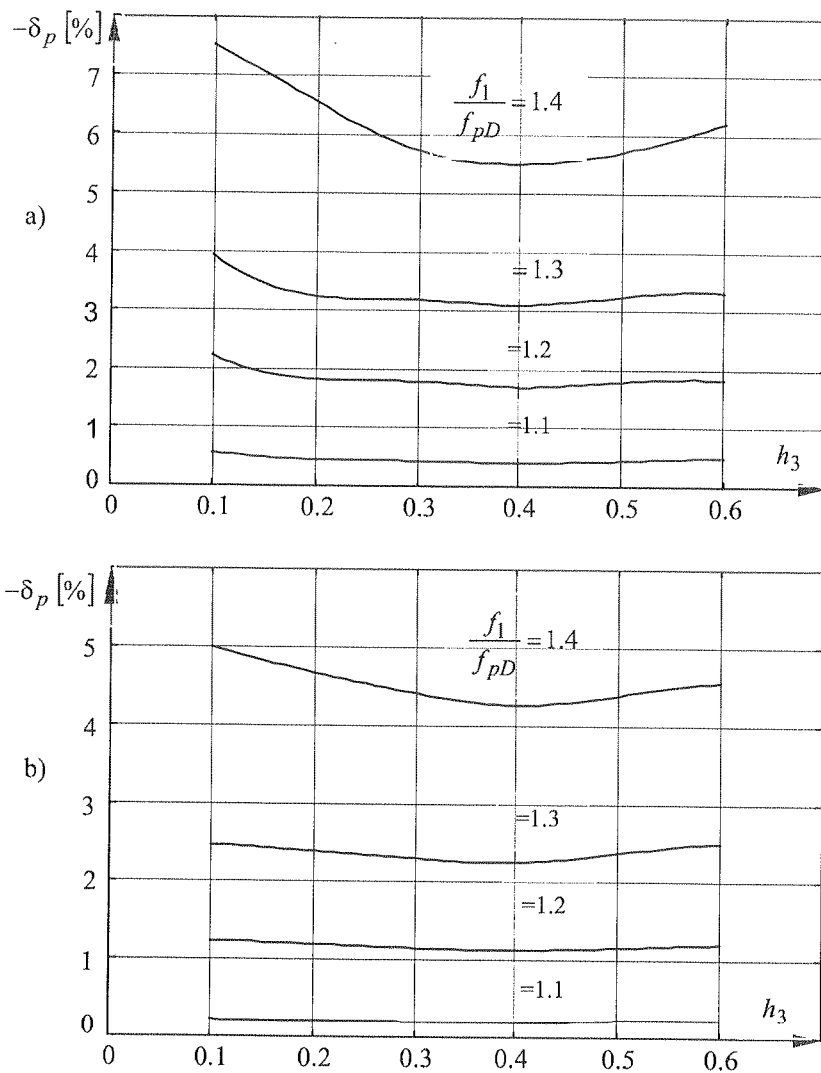
Z przebiegu charakterystyk przedstawionych na rys. 2 wynika, że jednoznaczny związek pomiędzy błędem przesterowania i udziałem h_3 trzeciej harmonicznej sygnału u_{t3} występuje, gdy wartość parametru h_3 jest mniejsza lub równa 0.4. Z tego względu do wyznaczania parametrów szybkościowych przetworników stosować można sygnał u_{t3} odkształcony trzecią harmoniczną, o udziale nie przekraczającym wartości granicznej $h_{3m} = 0.4$.

Górną graniczną wartość udziału harmonicznnych odkształcających sygnały u_{t23} (rys. 1b i c) wyznaczono na podstawie przebiegu charakterystyk z rys. 3 [11].

Rysunek ten przedstawia charakterystyki $\delta_p = f(h_2)$, wyznaczone numerycznie przy założeniu: $\frac{h_3}{h_2} = \text{const}$ i $\frac{f_1}{f_{pD}} = 1.4$ [11]. Harmoniczne sygnału u_{t23} mają największy

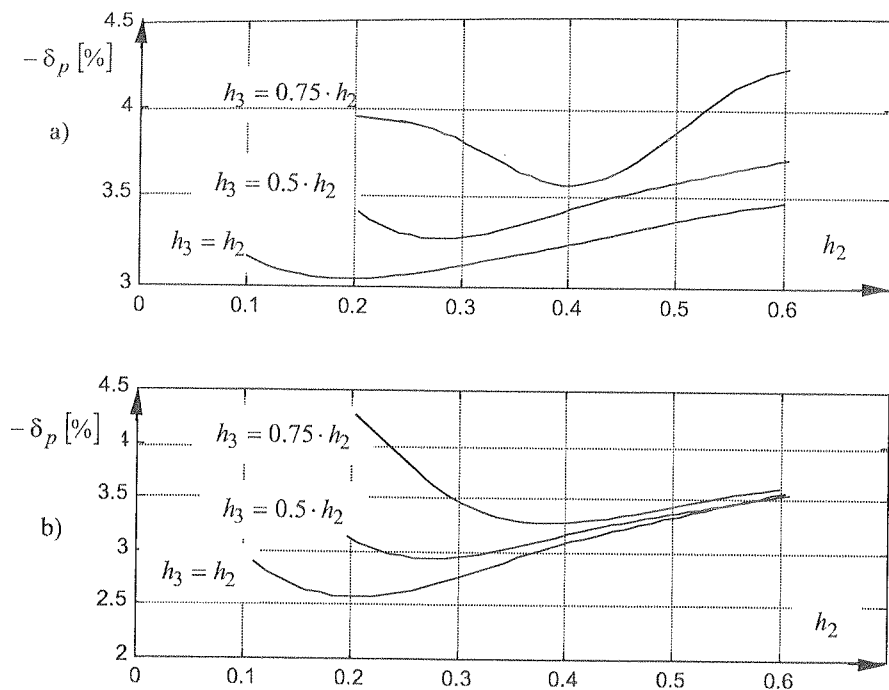
Rys. 2

Fig. 2. CI



Rys. 2. Charakterystyki $\delta_p = f(h_3)$ (dla $\frac{f_1}{f_{pD}}$) przetwornika wartości średniej (a) oraz przetwornika wartości skutecznej (b)

Fig. 2. Characteristics $\delta_p = f(h_3)$ (for $\frac{f_1}{f_{pD}}$) a mean value transducer (a) and a rms value transducer (b)



Rys. 3. Charakterystyki $\delta_p = f(h_2)$ (dla $\frac{f_1}{f_{pD}} = 1.4$ i $\frac{h_3}{h_2} = \text{const}$) przetwornika wartości średniej (a) oraz przetwornika wartości skutecznej (b)

Fig. 3. Characteristics $\delta_p = f(h_2)$ (for $\frac{f_1}{f_{pD}} = 1.4$ and $\frac{h_3}{h_2} = \text{const}$) a mean value transducer (a) and a rms value transducer (b)

wpływ na wartość błędu przesterowania, gdy udział tych harmoniczných spełnia zależność [11]:

$$h_3 = 0,75 \cdot h_2 \quad (2)$$

Z przebiegu charakterystyk $\delta_p = f(h_2)$ wyznaczonych dla $h_3 = 0,75 \cdot h_2$ wynika, że błąd przesterowania zależy jednoznacznie od udziału h_2 drugiej harmoniczných, gdy parametr h_2 jest mniejszy lub równy wartości granicznej $h_{2m} = 0.4$. Stąd po uwzględnieniu tej wartości z zał. (2) otrzymujemy górny graniczny udział trzeciej harmoniczných $h_{3m} = 0.3$ sygnałów u_{r23} .

W celu wyznaczenia dolnej granicznej wartości udziału harmoniczných sygnałów u_{r3} i u_{r23} określono przedział częstotliwości f_1 harmoniczných podstawowej tych sygnałów, w którym występuje, niezbędny do wyznaczania parametrów szybkościowych i błąd przesterowania przetwornika. Znajomość takiego przedziału częstotliwości f_1 pozwoli również na zweryfikowanie, wyznaczonych na podstawie przebiegu charakterystyk $\delta_p = f(h_k)$ (rys. 2 i 3), górnych wartości granicznych udziału ww. harmoniczných.

Przy wyznaczaniu błędu przesterowania (zal. 1) udział harmonicznych sygnałów u_{t3} i u_{t23} , z trzecią harmoniczną o zerowej fazie ($\varphi_3 = 0$ rad), powinien mieć wartość, przy której częstotliwość graniczna f_{pD} pasma przetwarzania mocy ww. sygnałów należy do zakresu niskich lub środkowych częstotliwości pasma przetwarzania przetwornika [11]. Spełnienie tego warunku pozwala na zminimalizowanie wpływu małosygnałowych błędów częstotliwościowych wejściowego stopnia przemiennoprądowego przetwornika na wartość błędu przesterowania [11, 13]. Częstotliwość graniczną f_{pD} pasma przetwarzania mocy odkształconych sygnałów wejściowych rozpatrywanych przetworników, wyznaczyć można z zależności [11, 13]:

$$f_{pD} = \frac{f_p}{\delta S_u} \quad (3)$$

W zal. (3) δS_u jest względną szybkością zmian sygnału odkształconego. Współczynnik δS_u zdefiniowano jako stosunek maksymalnej wartości modułu szybkości zmian δS_u tego sygnału do maksymalnej szybkości zmian S_{sm} sygnału sinusoidalnego, o takiej samej wartości skutecznej [11, 15]:

$$\delta S_u = \frac{\max|S_u|}{S_{sm}} \quad (4)$$

Z zal. (3) wynika, że wyznaczenie częstotliwości granicznej f_{pD} pasma przetwarzania mocy sygnałów u_{t3} i u_{t23} wymaga znajomości współczynników szybkości zmian $\delta S_{u_{t3}}$ i $\delta S_{u_{t23}}$ tych sygnałów. W tym celu należy określić również częstotliwość graniczną f_p pasma przetwarzania mocy sygnału sinusoidalnego, którego wartość równa jest wartości skutecznej sygnałów odkształconych u_{t3} i u_{t23} .

Zależności opisujące współczynniki szybkości zmian ww. sygnałów odkształconych mają postać następującą [11]:

— sygnału u_{t3} , gdy $\varphi_3 = 0$ rad:

$$\delta S_{u_{t3}} = \frac{1 + 3 \cdot h_3}{\sqrt{1 + h_3^2}} \quad (5)$$

— sygnałów u_{t23} , gdy $\varphi_3 = 0$ rad oraz $\varphi_2 = 0$ rad lub $\varphi_2 = \pi$ rad:

$$\delta S_{u_{t23}} = \frac{1 + 2 \cdot h_2 + 3 \cdot h_3}{\sqrt{1 + h_2^2 + h_3^2}} \quad (6)$$

Zal. (5) oraz zal. (6) uzyskano przekształcając odpowiednio zależność [11]:

$$\delta S_u = \frac{\max \left| \sum_{k=1}^n k \cdot h_k \cdot \cos(k\omega_1 t + \varphi_k) \right|}{\sqrt{1 + \sum_{k=2}^n h_k^2}} \quad (7)$$

opisującą współczynnik szybkości zmian sygnału odkształconego:

$$u = U_{m1} \sum_{k=1}^n h_k \sin(k\omega_1 t + \varphi_k) \quad (8)$$

gdzie: U_{m1} i h_k oznaczają odpowiednio amplitudę harmoniczną podstawowej oraz udział k -tej harmonicznego o amplitudzie U_{mk} : $h_k = \frac{U_{mk}}{U_{m1}}$, φ_k jest fazą k -tej harmonicznego.

Po przekształceniu zal. (3), z uwzględnieniem zal. (5) oraz zal. (6), otrzymujemy zależności opisujące względną wartość częstotliwości granicznej pasma przetwarzania mocy sygnałów testujących [11]:

— sygnału u_{r3} :

$$\frac{f_{pD}}{f_p} = \frac{\sqrt{1 + h_3^2}}{1 + 3 \cdot h_3} \quad (9)$$

— sygnałów u_{r23} :

$$\frac{f_{pD}}{f_p} = \frac{\sqrt{1 + h_2^2 + h_3^2}}{1 + 2 \cdot h_2 + 3 \cdot h_3} \quad (10)$$

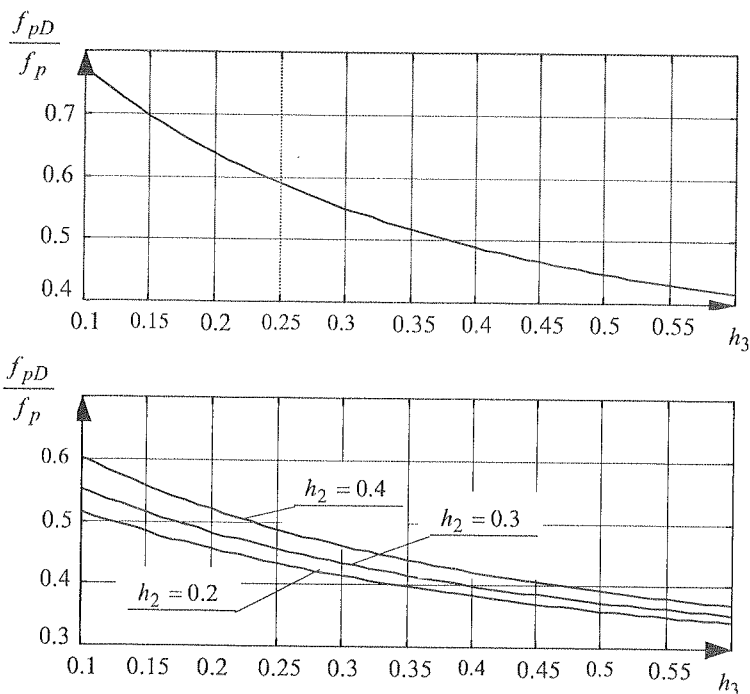
Ilustracją graficzną zal. (9) i zal. (10) są charakterystyki przedstawione na rys. 4.

Rys. 4a przedstawia charakterystykę $\frac{f_{pD}}{f_p} = f(h_3)$, wyznaczoną na podstawie zal. (9)

dla sygnału u_{r3} . Na rys. 4b przedstawiono charakterystyki $\frac{f_{pD}}{f_p} = f(h_3)$ ($h_2 = \text{const}$) dla sygnałów u_{r23} , do wyznaczenia których wykorzystano zal. (10).

Z charakterystyki przedstawionej na rys. 4a wynika, że względna częstotliwość graniczna $\frac{f_{pD}}{f_p}$ sygnału u_{r3} , którego trzecia harmoniczna ma udział $h_3 = 0.4$, równa jest 0.49. Zbliżoną wartość ma również, wyznaczona z charakterystyk z rys. 4b, względna częstotliwość graniczna $\frac{f_{pD}}{f_p}$ sygnałów u_{r23} , których druga i trzecia harmoniczna mają udział równy odpowiednio 0.4 i 0.3: $\frac{f_{pD}}{f_p} = 0.46$.

Przy założeniu, że częstotliwość graniczna f_p sygnału sinusoidalnego równa jest częstotliwości granicznej f_g pasma przetwarzania przyrządu, można stwierdzić, że uzyskane wartości względnej częstotliwości granicznej $\frac{f_{pD}}{f_p}$ sygnałów u_{r3} i u_{r23} , odkształconych harmonicznymi o udziale równym górnej wartości granicznej, należą do zakresu środkowych częstotliwości pasma przetwarzania przetwornika. Zakres ten ogranicza od dołu częstotliwość $0.3f_g$, zaś górna częstotliwość graniczna równa jest $0.7f_g$ [1, 2, 11] (f_g — częstotliwość graniczna pasma przetwarzania przetwornika). Do zakresu



Rys. 4. Charakterystyka $\frac{f_{pD}}{f_p} = f(h_3)$ dla sygnału u_{t3} (a) oraz charakterystyki $\frac{f_{pD}}{f_p} = f(h_3)$ ($h_2 = \text{const}$) dla sygnałów u_{t23} (b)

Fig. 4. Characteristics $\frac{f_{pD}}{f_p} = f(h_3)$ for the u_{t3} signal (a) and characteristics $\frac{f_{pD}}{f_p} = f(h_3)$ ($h_2 = \text{const}$) for u_{t23} signals (b)

środkowych częstotliwości pasma przetwarzania przetwornika należeć muszą również względne częstotliwości $\frac{f_{pD}}{f_p}$ sygnałów u_{t3} i u_{t23} odkształconych harmonicznymi, których udział ma dolną wartość graniczną. Udział ten wyznaczono zakładając, że częstotliwość względna $\frac{f_{pD}}{f_p}$ równa jest górnej częstotliwości granicznej zakresu środkowych częstotliwości pasma przetwarzania przetworników, tzn.: $\frac{f_{pD}}{f_p} = 0.7$. Z charakterystyki przedstawionej na rys. 4a wynika, że taką wartość częstotliwości granicznej $\frac{f_{pD}}{f_p}$ ma sygnał u_{t3} odkształcony trzecią harmoniczną o udziale $h_{3min} = 0.15$. Jest to dolna wartość graniczna udziału ww. harmoniczej. Postępując analogicznie z charakterystyk przedstawionych na rys. 4b, po uwzględnieniu zał. (2), wyznaczono dolną graniczną wartość udziału drugiej i trzeciej harmoniczej sygnałów u_{t23} : $h_{2min} = 0.25$, $h_{3min} = 0.19$.

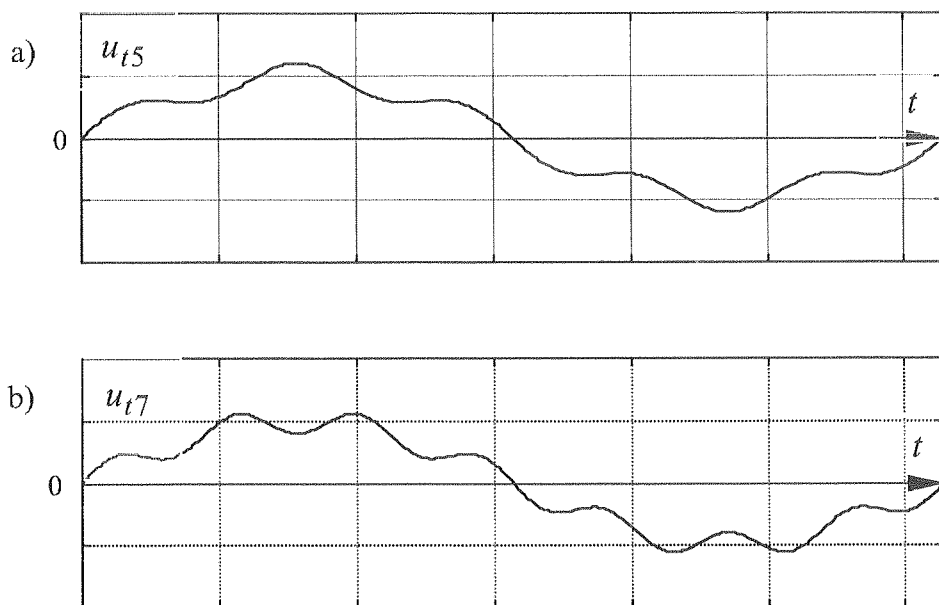
Reasumując dotychczasowe rozważania można stwierdzić, że parametry szybkościowe przetworników wyznaczyć można za pomocą sygnału u_{t3} (rys. 1a) odkształconego trzecią harmoniczną o udziale:

$$h_3 \in \langle 0.15, \dots, 0.4 \rangle \quad (11)$$

oraz za pomocą sygnałów u_{t23} odkształconych drugą i trzecią harmoniczną, których udziały spełniają warunek:

$$h_3 = 0.75 \cdot h_2 \text{ dla } h_3 \in \langle 0.1, \dots, 0.3 \rangle \text{ i } h_2 = \langle 0.25, \dots, 0.4 \rangle \quad (12)$$

Należy również wyjaśnić, dlaczego do wyznaczania parametrów szybkościowych przetworników nie należy stosować sygnału u_{t5} (rys. 5a) stanowiącego sumę pierwszej i piątej harmonicznej lub sygnału u_{t7} (rys. 5b) utworzonego przez sumę pierwszej i siódmej harmonicznej.



Rys. 5. Przebieg sygnałów u_{t5} (a) oraz u_{t7} (b) ($h_5 = h_7 = 0.2$, $\varphi_5 = \varphi_7 = 0$ rad)

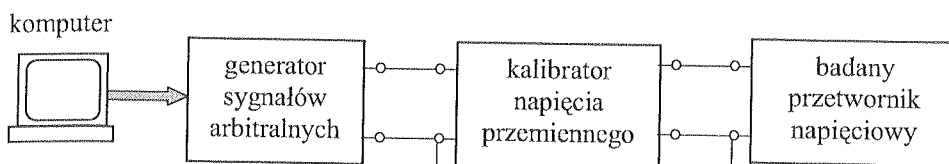
Fig. 5. Time-domain of u_{t5} (a) and u_{t7} (b) signals ($h_5 = h_7 = 0.2$, $\varphi_5 = \varphi_7 = 0$ rad)

Z rozważań autora przedstawionych w pracy [11] wynika, że zakres zmiany udziału h_5 i h_7 harmoniczných (piątej i siódmej), przy którym sygnały u_{t5} i u_{t7} mogą być przydatne do wyznaczania parametrów szybkościowych przetworników jest bardzo wąski tzn.: $h_5, h_7 \in \langle 0.1, 0.15 \rangle$. Występujące przy takim udziale ww. harmoniczných, wartości błędu przesterowania są czterokrotnie mniejsza, w stosunku do wartości błędu przesterowania przetworników na wejściu, których występują sygnały u_{t3} u_{t23} . W praktyce

oznacza to małą dokładność wyznaczania parametrów szybkościowych przetworników, za pomocą sygnałów odkształconych piątą i siódmą harmoniczną.

2.2. WYZNACZANIE PARAMETRÓW SZYBKOŚCIOWYCH

Schemat blokowy układu pomiarowego, za pomocą którego wyznaczyć można parametry szybkościowe przetworników napięciowych przedstawia rys. 6 [12].



Rys. 6. Schemat blokowy układu pomiarowego do wyznaczania częstotliwości f_{ph} przetworników napięciowych

Fig. 6. Block diagram of the measurement system for estimating of the f_{ph} frequency of voltage transducers

W układzie tym synteza przebiegu czasowego sygnałów testujących (u_{t3} lub u_{t23}) realizowana jest programowo w komputerze PC i następnie przebieg tych sygnałów w postaci cyfrowej przekazywany jest do pamięci generatora sygnałów arbitralnych. Sygnał wyjściowy tego generatora dołączony jest do wejścia kalibratora napięcia przemiennego, do którego zacisków roboczych podłączony jest badany przetwornik. W przedstawionym na rys. 6 układzie pomiarowym źródłem sygnałów testujących może być również odpowiednio połączone kalibratory napięcia [11].

W celu wyznaczenia parametru S_d przetwornika badanego wartość częstotliwość sygnałów testujących u_{t3} i u_{t23} , których harmoniczne składowe są zgodne w fazie, musi być ustalona na poziomie równym f_{pD} . Jest to graniczna częstotliwość pasma, w którym badany przetwornik przetwarza sygnały u_{t3} i u_{t23} z zerowym błędem przesterowania. Ze względu na sposób wyznaczania tego błędu (zal. 1) dokładne wyznaczenie zerowej jego wartości w praktyce nie jest możliwe. Spowodowana przesterowaniem szybkościowym minimalna wiarygodna różnica wskazań przetwornika badanego równa jest wartości błędu określonego wskaźnikiem jego klasy. Dlatego przyjęto, że parametr S_d wyznaczany jest przy częstotliwość f_{ph} ($f_{ph} > f_{pD}$) sygnałów testujących, które są przyczyną błędu przesterowania równego wartości błędu podstawowego przetwornika badanego [11, 12].

Sygnał testujący u_{t3} odkształcony trzecią harmoniczną (rys. 1a) ma symetryczny przebieg w półokresie oraz równe bezwzględne wartości maksymalnej szybkości narastania oraz maksymalnej szybkości opadania [11, 13]. Takie cechy sygnału u_{t3} pozwalają na wykorzystanie tego sygnału do badania przetworników do scharakteryzowania, których wystarczy informacja, o niezależnej od kierunku zmian sygnału wejściowego, dopuszczalnej szybkości S_d sygnału wejściowego [12]. Są to przetworni-

ki mierzące sygnały odkształcone, których harmoniczna podstawowa ma częstotliwość mniejszą lub równą 1 kHz [11]. Przy takich sygnałach wejściowych różnica wartości maksymalnej szybkości zmian S , występująca przy narastaniu i opadaniu sygnału wyjściowego, produkowanych obecnie analogowych układów i wzmacniaczy operacyjnych [3] ma pomijalnie mały wpływ na wartość błędu przesterowania przetworników wielkości elektrycznych. Wynika to z prac [6, 7, 10]. Wpływ różnicy wartości parametrów wzmacniaczy operacyjnych ujawnia się przy narastaniu i opadaniu sygnałów wejściowych o częstotliwości zbliżonej do górnej granicy pasma przetwarzania przetworników przeznaczonych do pracy z impulsowymi sygnałami wejściowymi. Dla takich przyrządów za pomocą sygnałów u_{t23} , odkształconych drugą i trzecią harmoniczną (rys. 1b i c), wyznaczyć należy częstotliwości graniczne f_{ph}^+ oraz f_{ph}^- [11]. Częstotliwości f_{ph}^+ wyznaczyć można za pomocą sygnału, którego harmoniczne składowe mają fazy zgodne (rys. 1b). Natomiast sygnał z rys. 1c pozwala na wyznaczenie częstotliwości f_{ph}^- . Przy wyznaczaniu parametrów szybkościowych ww. przyrządów przyjąć należy, że częstotliwość f_{ph} ma wartość określoną zależnością [11, 12]:

$$f_{ph} = \min [f_{ph}^+, f_{ph}^-] \quad (13)$$

Częstotliwość f_{ph} wyznaczyć można w sposób następujący [11, 12]. Za pomocą sygnału sinusoidalnego wyznaczyć należy przebieg częstotliwościowej charakterystyki $\delta_p = f(f_1)$ błędu badanego przyrządu. Następnie za pomocą odkształconego sygnału testującego (u_{t3} lub u_{t23}) wyznaczamy częstotliwościową charakterystykę $\delta_p = f(f_1)$ błędu przesterowania (zał. 1). Podczas badań wartość skuteczna sygnałów (u_{t3} i u_{t23}) musi być równa wartości skutecznej sygnału sinusoidalnego, za pomocą którego wyznaczono charakterystykę $\delta_p = f(f_1)$. Przy wyznaczaniu charakterystyki $\delta_p = f(f_1)$ zakres zmian częstotliwości f_1 ww. sygnałów od dołu ograniczony jest dolną częstotliwością graniczną f_d pasma przetwarzania badanego przyrządu. Ze względu na występowanie w sygnałach testujących trzeciej harmonicznej, górną granicę tego zakresu stanowi częstotliwość $f_g' = \frac{f_g}{3}$ (f_g jest górną częstotliwością graniczną pasma przetwarzania przetwornika badanego).

Przy wyznaczaniu charakterystyki $\delta_p = f(f_1)$ przyczyną błędu przesterowania badanego przyrządu jest zmiana fazy, z 0 rad na π rad, trzeciej harmonicznej sygnałów testujących. Konsekwencją zmiany fazy tej harmonicznej jest zmiana wartości średniej z modułu U_{lav} sygnałów testujących [11]. Przy wyznaczaniu charakterystyki $\delta_p = f(f_1)$ przetworników wartości średniej, niezależnie od fazy trzeciej harmonicznej, parametr U_{lav} sygnałów testujących powinien mieć stałą wartość. Z tego względu, przy zmianie fazy tej harmonicznej, amplituda harmonicznej podstawowej sygnału testującego (u_{t3} lub u_{t23}) musi zmieniać się zgodnie z zależnością [11, 12]:

$$U'_{m1} = U_{m1} \frac{1 + \frac{h_3}{3} \cos \varphi_3}{1 + \frac{h_3}{3} \cos (\varphi_3 + \Delta \varphi_3)} \quad (14)$$

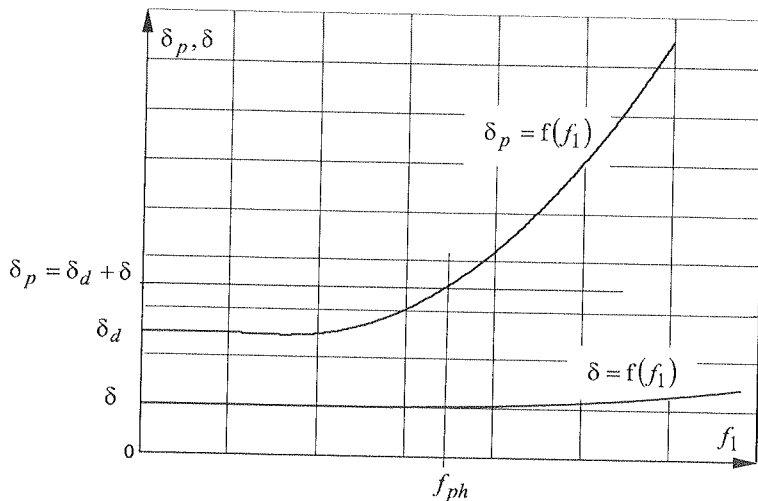
W zależności tej U_{m1} jest amplitudą harmoniczną podstawowej przed zmianą fazy φ_3 trzeciej harmonicznnej, U'_{m1} oznacza amplitudę ww. harmonicznnej po zmianie fazy tej harmonicznnej do wartości. $\varphi_3 + \Delta\varphi_3$. Udział h_3 trzeciej harmonicznnej musi być stały: tzn. $h_3 = \text{const}$, niezależnie od fazy trzeciej harmonicznnej.

Zmiana amplitudy harmonicznnej podstawowej z U_{m1} na U'_{m1} jest również przyczyną różnicy wartości szczytowej sygnału, którego harmoniczne składowe mają fazy zgodne oraz wartości szczytowej sygnału, którego trzecia harmoniczna jest w przeciwfazie w stosunku do harmonicznnej podstawowej. Ze względu na nieliniowość charakterystyki przetwarzania różnica wartości szczytowych sygnałów testujących może być przyczyną dodatkowego błędu przetwarzania badanego przetwornika nawet, gdy przetwornik ten nie jest przesterowany szybkościowo. Z tego względu, gdy harmoniczna podstawowa sygnałów testujących ma częstotliwość $f_1 = f_{ph}$, wartość błędu δ_p określona jest równaniem [11, 12]:

$$\delta_p = \delta_d + \delta \text{ dla } f_1 = f_{ph} \quad (15)$$

W zal. (15) δ jest błędem badanego przyrządu dla sygnału sinusoidalnego o częstotliwości f_{ph} , δ_d oznacza wartość błędu δ_p występującą przy sygnałach testujących o częstotliwości f_1 równej dolnej częstotliwości granicznej f_d pasma przetwarzania tego przyrządu.

Sposób wykorzystania przebiegu charakterystyk $\delta = f(f_1)$ oraz $\delta_p = f(f_1)$ do wyznaczania częstotliwości f_{ph} ilustruje rys. 7.



Rys. 7. Charakterystyki $\delta_p = f(f_1)$ i $\delta = f(f_1)$ (f_d — oznacza dolną częstotliwość graniczną pasma przetwarzania przetwornika)

Fig. 7. Characteristics $\delta_p = f(f_1)$ and $\delta = f(f_1)$ (f_d — signify low limit frequency of a transducers transmitted frequency band)

Aby częstotliwość f_{ph} można było wykorzystać do obliczenia parametru f_p badanego przetwornika należy sprawdzić, czy przyczyną występującego w zal. (15), przyrostu błędu δ_p ponad wartość δ_d , jest przesterowanie szybkościowe. Do tego celu wykorzystać można charakterystyki $\delta_p = f(f_1)$ wyznaczone sygnałem testującym z trzecią harmoniczną, której udział ma dwie różne wartości, przy zachowaniu stałej wartości skutecznej ww. sygnału. Jeżeli przyczyną przyrostu błędu δ_p , ponad wartość δ_d (rys. 7), jest przesterowanie szybkościowe, to pomiędzy częstotliwością graniczną f_{ph} sygnału testującego, którego trzecia harmoniczna ma udział h_3 i częstotliwością graniczną f_{ph} sygnału testującego, z trzecią harmoniczną o udziale $h_3 < h_3$, zachodzi zależność [11]

$$f'_{ph} \cong \frac{\delta S_u}{\delta S'_u} f_{ph} \quad (16)$$

W zal. (16) δS_u oraz $\delta S'_u$ oznaczają współczynniki szybkości zmian sygnału testującego odkształconego trzecią harmoniczną, o udziale odpowiednio równym h_3 i h_3 . Wartości współczynników δS_u oraz $\delta S'_u$ sygnałów wyznaczyć można za pomocą zal. (9) oraz zal. (10).

Dopuszczalna szybkość zmian S_d sygnału wejściowego badanego przetwornika równa jest maksymalnej szybkości zmian S_u sygnałów testujących (u_{t3} i u_{t23}), o częstotliwości f_{ph} , których trzecia harmoniczna ma fazę zgodną z fazą harmoniczej podstawowej [11, 12]:

$$S_d = S_u = 2 \cdot \pi \cdot f_{ph} \cdot U_{m1} \sum_{k=1, 2, 3} k \cdot h_k \quad (17)$$

gdzie: U_{m1} oznacza amplitudę harmoniczej podstawowej sygnału, h_k oznacza udział k -tej harmoniczej o amplitudzie $U_{mk} (h_k = \frac{U_{mk}}{U_{m1}})$.

Częstotliwość graniczną f_p wyznaczyć można przekształcając odpowiednio zal. (3) przy założeniu, że [11, 12]:

$$f_p = \delta S_u \cdot f_{ph} \quad (18)$$

gdzie: δS_u oznacza współczynnik szybkości zmian sygnału testującego u_{t3} lub u_{t23} , którego wartość opisują odpowiednio zal. (9) i (10).

Można wykazać, że w układzie pomiarowym z kalibratorem (rys. 6) częstotliwość graniczną f_{ph} można wyznaczyć z błędem wynikającym z klasy badanego przetwornika [11]. Z takim samym błędem wyznaczyć można, za pomocą zal. (17) i (18), parametry szybkościowe S_d i f_p tego przetwornika.

2.3. WYNIKI BADAŃ LABORATORYJNYCH

W celu sprawdzenia w praktyce pomiarowej, przedstawionego w artykule sposobu wyznaczania parametrów szybkościowych, przeprowadzono badania laboratoryjne

szerokopasmowych przyrządów elektronicznych. Badano multimetr firmy Datron typu 1281, multimetr firmy Fluke typu 8842A oraz przetwornik napięciowy firmy Analog Devices typu AD736. Za pomocą wymienionych przyrządów wyznaczyć można wartość skuteczną napięciowych sygnałów odkształconych o współczynnika szczytu nie większym od 3 w paśmie od 40Hz do 100kHz. Opis badań oraz ich wyniki opublikowano w pracach [11, 13].

Częstotliwościowe charakterystyki błędu przesterowania $\delta_p = f(f_1)$ badanych przyrządów wyznaczono za pomocą sygnału testującego u_{t3} (rys. 1 a) odkształconego trzecią harmoniczną o udziale 0,2 i 0,4, którego wartość skuteczna równa była 1V. Taką samą wartość skuteczną miał sygnał sinusoidalny, za pomocą którego wyznaczono częstotliwościowe charakterystyki błędu podstawowego $\delta_p = f(f_1)$.

Charakterystyki $\delta_p = f(f_1)$ i $\delta = f(f_1)$ wykorzystano do wyznaczenia częstotliwości δ_{ph} i za pomocą zał. (17) i (18) wyznaczono wartości parametrów S_d i f_p badanych przyrządów (tabela 1) [11].

Tabela 1

Parametry szybkościowe badanych przyrządów

Speed parameters of tested devices

	f_p [kHz]	S_d [V/ms]
multimetr Datron 1281	47.1	266.7
multimetr Fluke 8842A	43.9	248.8
przetwornik AD736	39.2	222.1

f_p - częstotliwość graniczna pasma przetwarzania mocy sygnałów sinusoidalnych,

S_d - dopuszczalna szybkość zmian sygnału wejściowego.

Analizując przedstawione w tabeli 1 wyniki obliczeń można stwierdzić, że parametr S_d badanych przyrządów przyjmuje stosunkowo małe wartości. Świadczy to o nieprzydatności tych przyrządów do pomiaru wartości skutecznej szybkozmiennych sygnałów impulsowych.

3. WYZNACZANIE PARAMETRÓW SZYBKOŚCIOWYCH PRZETWORNIKÓW MOCY

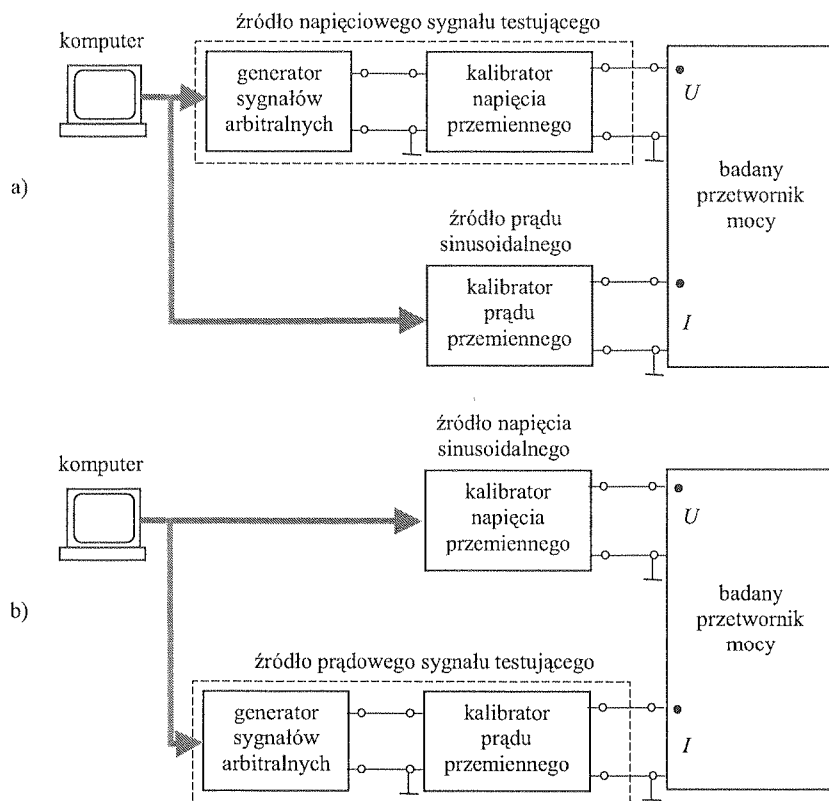
Parametry szybkościowe pomiarowych przetworników mocy wyznaczyć można, podobnie jak parametry szybkościowe przetworników napięciowych, za pomocą sygnałów testujących odkształconych trzecią harmoniczną (rys. 1a) lub za pomocą sygnałów testujących odkształconych drugą i trzecią harmoniczną (rys. 1b i c) [11, 12].

Dla przetworników przeznaczonych do pracy z sygnałami odkształconymi o częstotliwości przemysłowej (50–100 Hz) wyznaczyć należy jedynie graniczne szybkości zmian sygnałów wejściowych toru prądowego i napięciowego tzn. parametry S_{di}

oraz S_{du} . Do scharakteryzowania właściwości przetworników przeznaczonych do pracy z odkształconymi szybkozmiennymi sygnałami wejściowymi, których częstotliwość zmieniać się może w granicach określonych małosygnałowym pasmem przetwarzania przyrządu, stosować należy, oprócz wymienionych szybkości dopuszczalnych, również częstotliwości graniczne f_{pi} oraz f_{pu} pasma przetwarzania mocy sinusoidalnych sygnałów wejściowych: prądowego i napięciowego.

Sygnały testujące odkształcone drugą i trzecią harmoniczną (rys. 1b i c) stosować należy do badania przetworników mocy, dla których zachodzi konieczność wyznaczenia dopuszczalnej szybkości narastania S_d^+ oraz dopuszczalnej szybkości opadania S_d^- obu sygnałów wejściowych przetwornika.

Schematy blokowe układów pomiarowych, za pomocą których wyznaczyć można częstotliwość graniczną f_{ph} toru napięciowego oraz toru prądowego przetworników mocy, przedstawiają odpowiednio rys. 8a i b [12].



Rys. 8. Schemat blokowy układu pomiarowego do wyznaczania częstotliwości f_{ph} wejścia napięciowego (a) oraz wejścia prądowego (b) przetworników mocy

Fig. 8. Block diagram of a measurement system for estimating of the f_{ph} frequency of voltage input (a) and current input (b) of power transducers

W układach z rys. 8 odkształcone sygnały testujące (prądowy i napięciowy) wytwarzane są przez układy synchronizowanych zewnętrznie kalibratorów z dodatkowym wejściem. Do tego wejścia dołączony jest generator sygnałów arbitralnych, którego pamięć zawiera przebieg czasowy sygnału testującego.

Podczas badania przetworników mocy odkształcony sygnał testujący doprowadzić należy do wejścia badanego toru przyrządu. Sygnałem wejściowym drugiego toru jest sygnał sinusoidalny, o takiej samej częstotliwości. Ze względu na najbardziej niekorzystne warunki pracy przetwornika mocy czynnej, fazy obu sygnałów wejściowych powinny być zgodne. Przy badaniu przetworników mocy bierniej różnica faz obu sygnałów testujących (odkształconego i sinusoidalnego) musi być równa $\frac{\pi}{2}$ rad. Wartość skuteczna obu ww. sygnałów powinna być równa wartości skutecznej odpowiednich sygnałów sinusoidalnych, za pomocą których określano przebieg częstotliwościowej charakterystyki błędu podstawowego przetwornika. Sposób postępowania przy wyznaczaniu wartości częstotliwości granicznej f_{ph} badanego toru przetwornika mocy jest analogiczny, jak dla przetworników napięciowych. Uzasadnieniem przedstawionego sposobu badania przetworników mocy jest przebieg napięcia i prądu odbiorników zasilanych za pośrednictwem przekształtników energoelektronicznych [5]. Najczęściej tylko jeden z tych sygnałów ma przebieg silnie odkształcony, drugi zaś jest sygnałem o przebiegu zbliżonym do sinusoidy.

4. PODSUMOWANIE

Parametry szybkościowe (S_d i f_p) wykorzystać można przy doborze przetworników wielkości elektrycznych mierzących sygnały odkształcone. W artykule przedstawiono sposób wyznaczania tych parametrów.

Parametry szybkościowe wyznaczyć można pośrednio, gdy znana jest graniczna częstotliwość f_{ph} pasma przetwarzania mocy sygnału odkształconego o znanych i odpowiednio dobranych parametrach. Wartość częstotliwości granicznej f_{ph} przetworników napięciowych wyznaczyć można za pomocą sygnałów testujących i których maksymalną szybkość (narastania lub opadania) zmienić można przez odpowiednią zmianę kształtu sygnału. Do grupy takich sygnałów należy sygnał utworzony przez sumę pierwszej i trzeciej harmonicznej oraz sygnały, które tworzy suma pierwszej, drugiej i trzeciej harmonicznej. Przy wyznaczaniu częstotliwości f_{ph} przyjęto, że takie sygnały testujące przetwarzane są z błędem przesterowania o wartości równej wartości błędu podstawowego przetwornika badanego. Jeśli źródło sygnałów testujących ma odpowiednio dobrane parametry metrologiczne, to z takim błędem wyznaczyć można również parametry szybkościowe S_d i f_p badanego przyrządu. Niezbędną do występowania błędu przesterowania o wymaganej wartości zmianę maksymalnej szybkości zmian ww. sygnałów testujących uzyskać można zmieniając fazę trzeciej harmonicznej, z 0 rad na π rad. Wartość dopuszczalnej szybkości zmian sygnału wejściowego S_d badanego przyrządu równa jest maksymalnej szybkości zmian sygnału testującego, którego

trzecia harmoniczna ma fazę zgodną z fazą harmoniczej podstawowej. Częstotliwość graniczną f_p wyznaczyć można z zależność definiującą współczynnik szybkości zmian δS_u sygnałów testujących, z trzecią harmoniczną zerowej fazy.

Zaprezentowany w artykule sposób wyznaczania parametrów szybkościowych zastosować można również dla pomiarowych przetworników mocy. Dla tych przetworników wyznaczyć należy parametry szybkościowe obu torów wejściowych: prądowego i napięciowego. Przy wyznaczaniu tych parametrów sygnałem wejściowym jednego toru przetwornika jest odkształcony sygnał testujący, gdy w tym czasie, na wejściu drugiego toru występuje sygnał sinusoidalny.

5. BIBLIOGRAFIA

1. W. Golde, R. Śliwa: *Wzmacniacze operacyjne i ich zastosowania — podstawy teoretyczne*, Warszawa, WNT, 1982.
2. M. Herpy: *Analog Integrated Circuits*, Budapest, Akademiai Kiado, 1984.
3. Katalogi podzespołów firm: Analog-Devices, Burr-Brown oraz Harris 1990–2002.
4. Norma PN - EN 60688
5. M. Nowak, R. Barlik: *Poradnik inżyniera energoelektronika*, Warszawa, WNT, 1998.
6. K. Pacholski: *Błąd przesterowania przetworników napięciowych sygnałów impulsowych*, Kwartalnik Elektroniki i Telekomunikacji z. 1. 1995, ss. 69–87.
7. K. Pacholski: *The High-Speed Overloading Error of Voltage Transducers*, Measurement vol. 15, Nr 3, 1995, ss. 165–168.
8. K. Pacholski: *Wpływ przesterowania szybkościowego na właściwości metrologiczne przetworników mocy czynnej oraz przetworników napięciowych wartości średniej i skutecznej*, Zeszyty Naukowe Politechniki Łódzkiej Nr 759, 1996.
9. K. Pacholski, Z. Marks-Wojciechowska: *The Mathematical Model of Electronic Measuring Transducers for Processing Distorted Signals*, 9th International Symposium IMEKO 1997, Glasgow 1997, pp. 177–180.
10. K. Pacholski: *Wpływ przesterowania szybkościowego na właściwości metrologiczne przetworników pomiarowych mocy czynnej*, Materiały Konferencyjne EPN'97, Zielona Góra 1997, ss. 167–176.
11. K. Pacholski: *Przesterowanie szybkościowe i jego wpływ na właściwości metrologiczne przetworników pomiarowych wielkości elektrycznych*, Zeszyty Naukowe Politechniki Łódzkiej Nr 910, 2002.
12. K. Pacholski: *Sposób wyznaczania parametrów szybkościowych przetworników mocy oraz przetworników napięciowych wartości średniej i skutecznej*, Zgłoszenie Patentowe Nr P — 362336, 2003.
13. K. Pacholski: *Parametry szybkościowe przetworników pomiarowych wielkości elektrycznych*, Kwartalnik Elektroniki i Telekomunikacji z. 1. 2005, ss. 139–160.

K. PACHOLSKI

ESTIMATING OF SUITABILITY OF ELECTRIC QUANTITY TRANSDUCERS TO MEASUREMENT OF DEFORMED SIGNALS

Summary

This paper describes the way of estimating of speed parameters of the power transducer as well as average value and rms value transducer for measuring fast acting and deformed signals. These parameters

are: the acceptable rate of input signal S_d and the limit frequency f_p of the transmitted frequency band of the sinusoidal input signal. Recommended in the valid standard, method of estimating the metrological properties of mentioned transducer does not take into consideration parameters S_d and f_p . Values of these parameters could be estimated by using test signals with asymmetric fast acting runtime with maximum value, which could be changed by suitable change of signal shape. Examples of test signals are: sum of first and third harmonic or sum of first, second and third harmonic. To estimate speed parameters it is recommended to change test signal by change third harmonic from 0 to π radians.

Keywords: voltage transducers, power transducers, speed parameters, deformed signals

Wy
prze

n
P
w
w
co
o
z
m
za
m
tr
ci
Pr
pr
kl

St

Pro
jest szer

Wyznaczanie nieizotermicznych charakterystyk dławikowych przetwornic dc-dc o sterowaniu PWM w stanie ustalonym przy wykorzystaniu metody modeli uśrednionych

KRZYSZTOF GÓRECKI

*Katedra Elektroniki Morskiej, Akademia Morska w Gdyni,
ul. Morska 83, 81-225 Gdynia,
e-mail: gorecki@am.gdynia.pl*

Otrzymano 2005.09.01

Autoryzowano 2005.12.22

W pracy zaproponowano metodę wyznaczania charakterystyk przetwornic dc-dc w stanie ustalonym, przy uwzględnieniu samonagrzewania w elementach półprzewodnikowych. Proponowana metoda bazuje na koncepcji modeli uśrednionych i pozwala na wyznaczenie wartości napięcia wyjściowego i sprawności, a także, istotnych z punktu widzenia niezawodności przetwornicy, temperatur wnętrza zawartych w niej tranzystorów i diod. W opracowanej metodzie, realizowanej przy wykorzystaniu programu SPICE, formułuje się cztery obwody. Pierwszy z nich stanowi sieć złożoną ze sterowanych źródeł napięciowych oraz rezystorów, której struktura wynika z opisanej w literaturze zasady formułowania uśrednionych modeli przetwornic. Drugi obwód zawiera szeregowo połączone tranzystor i diodę, opisane za pomocą modeli wbudowanych w programie SPICE oraz sterowane źródła napięciowe, modelujące wpływ temperatury na napięcie między zaciskami wyjściowymi włączonego tranzystora oraz spadek napięcia na diodzie spolaryzowanej w kierunku przewodzenia. Trzeci i czwarty obwód stanowią analogi elektryczne modeli termicznych tranzystora i diody. Przedstawiono sposób formułowania wymienionych obwodów na przykładzie dławikowych przetwornic buck oraz boost. Poprawność opracowanej metody zweryfikowano za pomocą klasycznej analizy stanów przejściowych, prowadzonej aż do uzyskania stanu ustalonego.

Słowa kluczowe: przetwornice dc-dc, analiza elektrotermiczna, SPICE, modele uśrednione

1. WPROWADZENIE

Problem komputerowej analizy układów impulsowych, w tym przetwornic dc-dc, jest szeroko rozważany w literaturze [1, 2, 3, 4, 5, 6, 7, 8, 9, 10]. Jak wynika m.in. z prac

[2, 11], komputerowa analiza układów impulsowego przetwarzania energii jest realizowana na różnych poziomach szczegółowości opisu zjawisk zachodzących w badanym układzie. Z punktu widzenia użytkownika bardzo ważne są charakterystyki przetwornicy w stanie ustalonym, które określają zarówno zakres możliwych do uzyskania zmian napięcia wyjściowego przetwornicy, jak również jej sprawności energetycznej.

Jedną z metod wyznaczania charakterystyk przetwornic dc-dc o sterowaniu PWM w stanie ustalonym jest metoda wykorzystująca uśrednione modele tych układów. Metoda ta została opisana m.in. w pracach [7, 9, 10, 12]. W metodzie tej są formułowane schematy zastępcze badanej przetwornicy przy założeniu, że napięcia na elementach półprzewodnikowych mają przebieg prostokątny, a czasy trwania przełączeń są pomijalnie krótkie. Przy formułowaniu uśrednionych modeli przetwornic elementy półprzewodnikowe są modelowane za pomocą rezystorów o skokowo zmienianej wartości rezystancji między dwoma ustalonymi wartościami - rezystancją włączenia R_{ON} oraz rezystancją wyłączenia R_{OFF} .

Wyróżnić można dwie odmiany tej metody. W pierwszej z nich dioda i tranzystor modelowane są identycznie ($R_{ON} = 0, R_{OFF} = \infty$), np. [10, 13, 14, 15]. Z kolei, w drugiej odmianie rozważanej metody [7, 9, 10, 12] tranzystor modelowany jest za pomocą klucza ($R_{ON} > 0, R_{OFF} = \infty$), natomiast dioda — za pomocą szeregowego połączenia klucza idealnego ($R_{ON} = 0, R_{OFF} = \infty$) ze źródłem napięcia o stałej wydajności V_D oraz z rezystorem o rezystancji R_D . Stosując rozważaną metodę otrzymuje się wzór analityczny, opisujący zależność napięcia wyjściowego U_{wy} przetwornicy od współczynnika wypełnienia d sygnału sterującego, np. [10, 13, 14, 15]. W przypadku stosowania drugiej odmiany rozważanej metody uzyskuje się wzory analityczne, opisujące zależność napięcia wyjściowego przetwornicy i jej sprawności od współczynnika wypełnienia sygnału sterującego, rezystancji obciążenia, napięcia wejściowego, rezystancji R_{ON} i R_D oraz napięcia V_D .

W pracach [7, 9, 11] zaproponowano sposób implementacji przedstawionych metod modeli uśrednionych w programie SPICE. Uzyskanie charakterystyk przetwornic w stanie ustalonym wymaga wówczas przeprowadzenia analizy stałoprądowej (dc).

Niestety znane z literatury odmiany metody modeli uśrednionych nie uwzględniają wpływu poziomów napięcia sterującego tranzystor kluczujący ani nieliniowości charakterystyk elementów półprzewodnikowych na charakterystyki przetwornic w stanie ustalonym.

W pracy [16] zaproponowano metodę przyspieszonego obliczania charakterystyk przetwornic o sterowaniu PWM w stanie ustalonym. W metodzie tej iteracyjnie wykonywane są analizy stanów przejściowych przy warunku początkowym uzyskanym z modelu uśrednionego, a z kolei wartości spadków napięć na zaciskach diody i tranzystora w modelu uśrednionym są uzyskiwane z analizy stanów przejściowych. Czas trwania obliczeń metodą z pracy [16] jest kilkukrotnie dłuższy od czasu trwania obliczeń wykonywanych przy wykorzystaniu metody z prac [7, 9, 10].

W omówionych powyżej literaturowych metodach analizy pominięto wpływ ważnego zjawiska fizycznego — efektu samonagrzewania [17] na charakterystyki prze-

twor
cie e
tego
temp
rakter
rakter
troter
elektr
sko sa
a war
grzew
Dodat
układu
W
klasy
metod
tod lite
półprz
peratur
sposób
ost ora
analizy

Jak
koncept
sji meto
element
szerego
propono
przez st
rozważa
Zgo
nie stało
rezystor
tyczne, j
prądów
Real
cztery pi
1. Opra
odpo

twornicy. Samonagrzewanie jest skutkiem wydzielania energii elektrycznej w elemencie elektronicznym i zamiany jej na ciepło przy nieidealnych warunkach chłodzenia tego elementu. W wyniku tego zjawiska wzrasta temperatura wnętrza elementu ponad temperaturę otoczenia oraz zmienia się przebieg jego charakterystyk, nazywanych charakterystykami nieizotermicznymi [17]. W celu wyznaczenia nieizotermicznych charakterystyk elementów i układów elektronicznych należy przeprowadzić analizę elektrotermiczną [17], w której wykorzystywane są elektrotermiczne modele elementów elektronicznych, uwzględniające samonagrzewanie. Jak wynika z prac [18, 19] zjawisko samonagrzewania w istotny sposób wpływa na charakterystyki przetwornic dc-dc, a wartości napięcia wyjściowego przetwornicy obliczone przy pominięciu samonagrzewania mogą być nawet o 30% zawyżone w stosunku do wartości rzeczywistych. Dodatkowo, wzrost temperatur wnętrza diody i tranzystora ogranicza niezawodność układu, a nawet może spowodować uszkodzenie przetwornicy.

W pracy zaproponowano nową metodę wyznaczania charakterystyk rozważanej klasy przetwornic dc-dc w stanie ustalonym. Proponowana metoda należy do grupy metod wykorzystujących uśrednione modele przetwornic i w przeciwieństwie do metod literaturowych, umożliwia uwzględnienie wpływu samonagrzewania w elementach półprzewodnikowych na rozważane charakterystyki oraz umożliwia wyznaczenie temperatur wnętrza tych elementów. W kolejnych rozdziałach opisano ideę nowej metody, sposób formułowania uśrednionych elektrotermicznych modeli przetwornic buck i boost oraz wyniki weryfikacji poprawności metody w oparciu o wyniki elektrotermicznej analizy stanów przejściowych wymienionych przetwornic.

2. IDEA METODY

Jak zaznaczono we Wprowadzeniu, proponowana metoda analizy wykorzystuje koncepcję uśrednionych modeli przetwornic dc-dc. Jednak, o ile w klasycznej wersji metody modeli uśrednionych [7, 9, 10] w schemacie zastępczym przetwornicy, elementy półprzewodnikowe zastępowane są przez rezystory liniowe (tranzystor) oraz szeregowo połączenie źródła napięciowego o stałej wydajności i rezystora (dioda), to w proponowanej metodzie oba wymienione elementy półprzewodnikowe są zastępowane przez sterowane źródła napięciowe, których wydajności zależą od parametrów modeli rozważanych elementów, a także od parametrów sygnału sterującego tranzystor.

Zgodnie z koncepcją metody modeli uśrednionych [10] niezbędne jest stworzenie stałoprądowego analogu obwodowego badanej przetwornicy, zawierającego jedynie rezystory i sterowane źródła napięciowe. W analogu tym muszą być zachowane identyczne, jak w rozważanej przetwornicy, relacje między wartościami średnimi napięć i prądów zaciskowych.

Realizacja proponowanej metody wymaga wykonania siedmiu kroków, przy czym cztery pierwsze są takie same, jak w klasycznej metodzie modeli uśrednionych [10]:

1. Opracowanie trzech obwodów zastępczych badanej przetwornicy, z których jeden odpowiada stanowi włączenia tranzystora, drugi — stanowi wyłączenia tranzystora

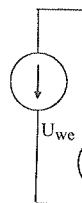
i włączenia diody, natomiast trzeci — stanowi wyłączenia tranzystora i wyłączenia diody. W stanie włączenia tranzystor jest reprezentowany przez źródło napięciowe U_{ON} , a w stanie wyłączenia — przez rozwarcie. Dioda jest reprezentowana w zakresie zaporowym przez rozwarcie, a w zakresie przewodzenia przez źródło napięciowe o wydajności U_D . Jak wykazano w pracach [20, 21], pominięcie prądów wstecznych elementów półprzewodnikowych, nie powoduje zauważalnych zmian w przebiegu charakterystyk zaciskowych przetwornic dławikowych, a w sposób istotny upraszcza analizę tych układów.

2. Sformułowanie równań opisujących napięcie na dławiku i prąd kondensatora dla sformułowanych w punkcie 1 obwodów zastępczych przy założeniu, że składowe zmienne napięć i prądów wynikające z przełączania są pomijalnie małe w stosunku do ich składowych stałych.
3. Sformułowanie i przyrównanie do zera wyrażeń opisujących wartość średnią napięcia na dławiku przy wykorzystaniu zasady ciągłości strumienia w cewce oraz opisujących wartość średnią prądu kondensatora, korzystając z zasady zachowania ładunku w kondensatorze.
4. Sformułowanie modelu obwodowego przetwornicy opartego na otrzymanych w punkcie 3 równaniach, opisujących zachowanie się przetwornicy w stanie ustalonym przy uwzględnieniu strat w elementach półprzewodnikowych i w dławiku.
5. Opracowanie obwodu złożonego z szeregowego połączenia tranzystora kluczującego, diody, dwóch sterowanych źródeł napięciowych modelujących zmiany napięcia wyjściowego tranzystora oraz napięcia przewodzenia diody, wywołane zmianą temperatury wnętrza tych elementów, źródła napięciowego o wydajności równej wysokiemu poziomowi napięcia sterującego tranzystor kluczujący, połączonego przez rezystor z elektrodą sterującą tego tranzystora oraz ze sterowanego źródła prądowego, którego wydajność równa jest średniej wartości prądu dławika w stanie ustalonym. Napięcia na zaciskach wyjściowych tranzystora i diody równe są wartościom napięć U_{ON} oraz U_D , występujących w opisie wydajności źródeł sterowanych zawartych w obwodzie głównym przetwornicy, sformułowanym w punkcie 4.
6. Sformułowanie stałoprądowych modeli termicznych tranzystora i diody, opisujących zależność temperatur wnętrza od wydzielanej w nich mocy, w postaci obwodowej. W rozważanej analizie może być zastosowany stałoprądowy model termiczny ze względu na duże wartości termicznych stałych czasowych w modelu przejściowej impedancji termicznej elementów półprzewodnikowych w porównaniu z okresem sygnału sterującego tranzystor. Wówczas, jak wynika z pracy [22] wartość temperatury wnętrza elementu w stanie ustalonym przy jego pracy impulsowej stanowi sumę temperatury otoczenia oraz iloczynu rezystancji termicznej tego elementu i wartości średniej wydzielanej w nim mocy.
7. Sformułowanie wzorów opisujących czas t_p , w którym płynie prąd dławika przy nieciąglym trybie pracy przetwornicy oraz wzorów opisujących średnie wartości prądu diody oraz prądu głównego tranzystora.

Pr
lacji c
klucz
sposób
przetw
źródł

UL
boost -
są pom
dło nap
przetw
w zakr
kierunk
trybie p
łączani
pomini
ogranic
model :

a)

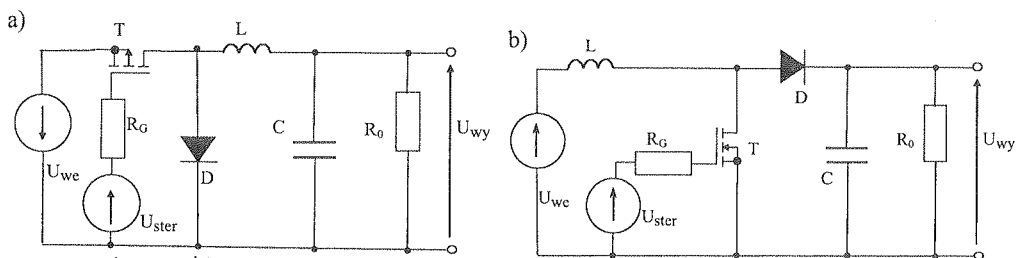


W
przetwor
pokazan
— 4 um
czania sp
rys. 2b i
6 algory
wyłączer

3. ZASTOSOWANIE NOWEJ METODY DO ANALIZY PRZETWORNIC DŁAWIKOWYCH

Przedstawioną w poprzednim rozdziale metodę analizy wykorzystano do symulacji charakterystyk przetwornic dławikowych boost i buck, w których elementem kluczującym jest tranzystor mocy MOS. W dalszej części rozdziału przedstawiono sposób formułowania obwodów, występujących w układzie zastępczym rozważanych przetwornic oraz uzasadniono postać wzorów opisujących wydajności odpowiednich źródeł sterowanych, zawartych w tych obwodach.

Układ przetwornicy buck zamieszczono na rys.1a, natomiast układ przetwornicy boost - na rys.1b. Przyjmując, że czasy przełączania elementów półprzewodnikowych są pomijalnie krótkie w relacji do okresu sygnału sterującego, wytwarzanego przez źródło napięciowe U_{ster} , można wyróżnić w każdym okresie tego sygnału trzy fazy pracy przetwornicy. W pierwszej z nich prąd dławika L płynie przez tranzystor T , pracujący w zakresie liniowym. W drugiej fazie prąd dławika płynie przez spolaryzowaną w kierunku przewodzenia diodę D . W trzeciej fazie, która występuje tylko w nieciągłym trybie pracy przetwornicy, prąd dławika nie płynie. Ze względu na krótkie czasy przełączania elementów półprzewodnikowych w relacji do okresu sygnału kluczującego, pominięto w analizie zjawiska zachodzące w czasie przełączania tych elementów, co ogranicza od góry zakres częstotliwości sygnału sterującego, dla których rozważany model zapewnia uzyskanie poprawnych wyników obliczeń.

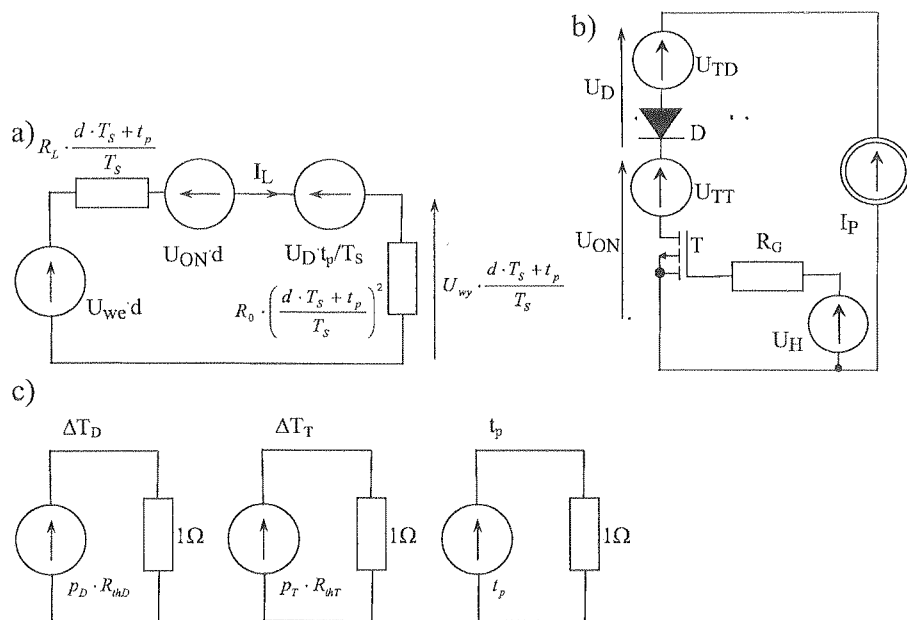


Rys. 1. Schemat dławikowej przetwornicy buck (a) oraz boost (b)

Fig. 1. The diagram of the buck converter (a) and the boost converter (b)

W wyniku realizacji algorytmu z rozdziału 2 uzyskuje się schemat zastępczy przetwornicy buck, pokazany na rys. 2 oraz schemat zastępczy przetwornicy boost pokazany na rys. 3. Główny obwód przetwornicy, wynikający z realizacji punktów 1 — 4 umieszczono odpowiednio na rys. 2a i rys. 3a, obwód pomocniczy do wyznaczenia spadku napięcia na diodzie i tranzystorze (punkt 5 algorytmu) umieszczono na rys. 2b i rys. 3b, natomiast obwodowe modele termiczne tranzystora i diody (punkt 6 algorytmu) oraz obwód do wyznaczenia czasu t_p , w którym płynie prąd cewki po wyłączeniu tranzystora (punkt 7 algorytmu) zamieszczono na rys. 2c i rys. 3c.

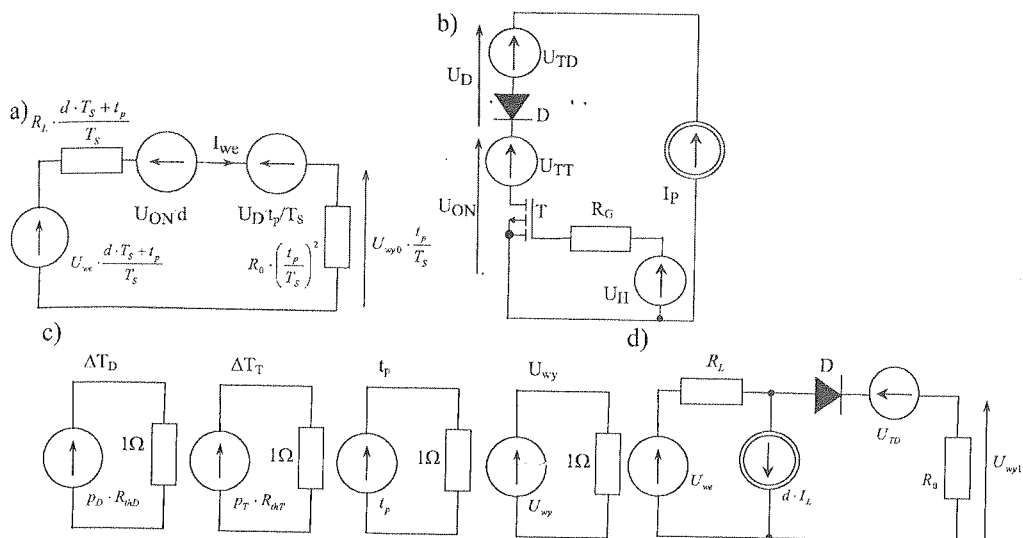
Występujące w opisie wydajności źródeł napięciowych na rys. 2a i rys.3a, spadki napięcia U_{ON} na włączonym tranzystorze oraz U_D na diodzie spolaryzowanej w kierunku przewodzenia są wyznaczane, w pokazanym na rys. 2b oraz rys. 3b, obwodzie pomocniczym. Przetwornica boost pracuje impulsowo tylko w przypadku, gdy napięcie U_{ON} na załączonym tranzystorze jest mniejsze od napięcia wejściowego U_{we} . Jeżeli ten warunek nie jest spełniony, to dioda D przewodzi w sposób ciągły i napięcie wyjściowe jest równe różnicy napięcia wejściowego i spadku napięcia na diodzie spolaryzowanej w kierunku przewodzenia. Wówczas obwód z rys. 3a powinien być zastąpiony przez obwód z rys. 3d.



Rys. 2. Ekwiwalent obwodowy elektrotermicznego modelu uśrednionego przetwornicy buck

Fig. 2. The network representation of the electrothermal average model of the buck converter

Na rys. 2 i 3, symbol d oznacza współczynnik wypełnienia impulsów sterujących tranzystor, T_s jest okresem powtarzania tych impulsów, natomiast U_H wartością poziomu wysokiego impulsów, R_L jest rezystancją szeregową cewki L , przez którą płynie prąd o wartości średniej I_L . U_{ON} oznacza spadek napięcia dren-źródło na tranzystorze T pracującym w zakresie liniowym, natomiast U_D — spadek napięcia na diodzie spolaryzowanej w kierunku przewodzenia. Źródła napięciowe U_{TT} oraz U_{TD} reprezentują zmiany napięcia na tranzystorze i diodzie spowodowane wzrostem temperatur ich wnętrza ponad temperaturę otoczenia, odpowiednio ΔT_T oraz ΔT_D . Moce wydzielane w diodzie i w tranzystorze wynoszą odpowiednio p_D oraz p_T , natomiast rezystancje termiczne diody i tranzystora oznaczono symbolami R_{thD} oraz R_{thT} . Prąd I_P jest wartością średnią prądów diody i tranzystora w czasie ich przewodzenia.



Rys. 3. Ekwiwalent obwodowy elektrotermicznego modelu uśrednionego przetwornicy boost

Fig. 3. The network representation of the electrothermal average model of the boost converter

Obwody z rys. 2a i 3a utworzono stosując metodę opisaną szczegółowo w pracy [16], przy czym uwzględniono dodatkowo możliwość pracy przetwornicy w trybie nieciągłym. Łatwo można zauważyć, że w przypadku pracy przetwornicy w trybie ciągłym (wówczas $t_p = (1 - d)T_S$) oraz przy pominięciu rezystancji cewki R_L , obwody te zredukują się do postaci przedstawionej w pracy [16].

Konstruując układ z rys. 2b i 3b wykorzystano spostrzeżenie, że średnia wartość prądu diody oraz tranzystora w czasie ich przewodzenia jest jednakowa, ponieważ wartość chwilowa tych prądów zmienia się liniowo w czasie przewodzenia wymienionych elementów oraz przyjmuje identyczne wartości maksymalne i minimalne. Przy pracy w trybie ciągłym wartości średnie prądów diody, tranzystora i cewki są jednakowe, natomiast przy pracy w trybie nieciągłym prąd I_p dany jest wzorem

$$I_p = I_L \cdot \frac{T_S}{d \cdot T_S + t_p} \quad (1)$$

Wydajności źródeł napięciowych modelujących temperaturowe zmiany napięcia przewodzenia diody oraz tranzystora uzyskano wykorzystując koncepcję hybrydowych liniowych modeli elektrotermicznych elementów półprzewodnikowych [17, 23, 24]. Wydajności rozważanych źródeł dane są wzorami

$$U_{TT} = r_{T0} \cdot I_L \cdot \frac{T_S}{d \cdot T_S + t_p} \cdot \alpha_{TT} \cdot \Delta T_T \quad (2)$$

$$U_{TD} = \alpha_{UD} \cdot \Delta T_D + r_{SD} \cdot I_L \cdot \frac{T_S}{d \cdot T_S + t_p} \cdot \alpha_{TD} \cdot \Delta T_D \quad (3)$$

gdzie r_{T0} oznacza rezystancję włączenia tranzystora MOS w temperaturze otoczenia, α_{TT} jest temperaturowym współczynnikiem zmian napięcia U_{ON} , α_{UD} — temperaturowym współczynnikiem zmian napięcia przewodzenia diody, r_{SD} — rezystancją szeregową diody, α_{TD} — temperaturowym współczynnikiem zmian rezystancji szeregowej diody.

Ponieważ w przetwornicach dc-dc zawsze okres sygnału sterującego tranzystor kluczujący jest znacznie krótszy od najkrótszej termicznej stałej czasowej, wartość średnia tej nadwyżki równa jest praktycznie iloczynowi wartości średniej mocy wydzielanej w tych elementach i ich rezystancji termicznej [22]. Wartości średnie mocy równe są iloczynowi prądu tych elementów, płynącego w obwodzie z rys. 2b lub 3b, przez napięcie na ich zaciskach oraz iloraz czasu przepływu prądu przez element w czasie jednego okresu sygnału sterującego do tego okresu. W przypadku diody wartość średnia wydzielanej w niej mocy dana jest wzorem

$$p_D = U_D \cdot I_L \cdot \frac{t_p}{d \cdot T_S + t_p} \quad (4)$$

natomiast wartość średnia mocy wydzielanej w tranzystorze wynosi

$$p_T = U_{ON} \cdot I_L \cdot \frac{d \cdot T_S}{d \cdot T_S + t_p} \quad (5)$$

W przypadku pracy w trybie ciągłym czas $t_p = (1 - d)T_S$, natomiast w trybie nieciągłym wartość czasu t_p wyznaczono przez przyrównanie średniego prądu obciążenia i średniego prądu dławika. Uzyskano następującą zależność opisującą czas t_p dla przetwornicy buck

$$t_p = \frac{-d \cdot T_S + \sqrt{d^2 \cdot T_S^2 + 8 \cdot \frac{U_{wy} \cdot T_S \cdot L}{R_0 \cdot (U_{wy} - U_D)}}}{2} \quad (6)$$

oraz dla przetwornicy boost

$$t_p = 2 \cdot \frac{U_{wy} \cdot L}{R_0 \cdot (U_{wy} - U_D) \cdot d} \quad (7)$$

Wzory (6) i (7) są słuszne tylko przy pracy w trybie nieciągłym. Jeżeli wartość uzyskana z tych wzorów jest większa niż $(1 - d)T_S$, to oznacza, że przetwornica pracuje w trybie ciągłym i wówczas $t_p = (1 - d)T_S$.

Jak widać, przy wykorzystaniu przedstawionej metody, uzyskuje się charakterystyki przetwornicy w stanie ustalonym za pomocą analizy stałoprądowej. W celu zabezpieczenia przed nadmiarem zmiennoprzecinkowym zastosowano w opisie niektórych źródeł funkcji limit. Ponieważ napięcie na rezystancji obciążenia R_0 nie jest równe napięciu wyjściowemu przetwornicy, dlatego zastosowano dodatkowe źródło napięciowe

E_{wy} , na którego zaciskach napięcie równe jest napięciu wyjściowemu przetwornicy. Napięcie wyjściowe U_{wy} przetwornicy boost jest dane wzorem

$$U_{wy} = \begin{cases} U_{wy0} & \text{gdy } U_{wy0} > U_{wy1} \\ U_{wy1} & \text{gdy } U_{wy0} < U_{wy1} \end{cases} \quad (8)$$

gdzie napięcie U_{wy0} uzyskano w układzie z rys. 3a, natomiast napięcie U_{wy1} w układzie z rys. 3c.

Z kolei, sprawność energetyczna przetwornicy buck wyrażona jest wzorem

$$\eta = \frac{U_{wy}^2 \cdot (d \cdot T_S + t_p)}{R_0 \cdot U_{we} \cdot d \cdot I_L \cdot T_S} \quad (9)$$

natomiast dla przetwornicy boost opisuje ją zależność

$$\eta = \frac{U_{wy}^2}{R_0 \cdot U_{we} \cdot I_{we}} \quad (10)$$

Opracowane modele przetwornic buck oraz boost zaimplementowano w programie SPICE w postaci podukładów.

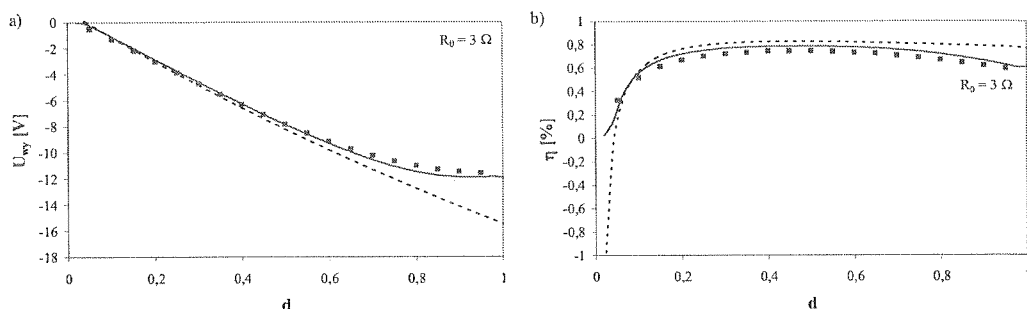
4. WERYFIKACJA POPRAWNOŚCI METODY

W celu zilustrowania poprawności i przydatności zaproponowanej metody przeprowadzono analizy przetwornic buck oraz boost przy wykorzystaniu elektrotermicznej analizy stanów przejściowych, opisaną w pracy [25, 26], analizy stałoprądowej przy wykorzystaniu zaproponowanej metody oraz analizy stałoprądowej przy wykorzystaniu metody modeli uśrednionych z pracy [10]. W analizach wykonanych za pomocą programu SPICE, przyjęto następujące wartości elementów biernych: dla przetwornicy buck: $U_{we} = -12$ V, $L = 100 \mu\text{H}$, $R_L = 0,2 \Omega$, $C = 470 \mu\text{F}$ oraz dla przetwornicy boost: $U_{we} = 20$ V, $L = 50 \mu\text{H}$, $R_L = 0,2 \Omega$, $C = 470 \mu\text{F}$. W analizach stanów przejściowych oraz w nowej metodzie wykorzystano następujące wartości parametrów modelu tranzystora MOS: Level = 3, Kappa = 0.2, Tox = 100 nm, $U_0 = 600 \text{ cm}^2 \text{V}^{-1} \text{s}^{-1}$, Phi = 0.6 V, RS = 6.382 m Ω , Kp = 20.85 $\mu\text{A/V}^2$, W = 0.68 m, L = 2 μm , Vto = 3.879 V, Rd = 0.6703 Ω , TT = 710 ns, Rds = 2.222 M Ω , Cbd = 1.415 nF, Pb = 0.8 V, Mj = 0.5, Fc = 0.5, Cgso = 1.625 nF/m, Cgdo = 133.4 pF/m, Rg = 0.6038 Ω , Is = 56.03 pA oraz parametrów modelu diody: Is = 53.41 pA, N = 1.185, RS = 0.12 Ω , trs1 = $3 \cdot 10^{-3} \text{ K}^{-1}$, Ikf = 3.501 mA, Cjo = 325 pF, M = 0.3333, Vj = 0.75 V, Fc = 0.5, Isr = 100 pA, Nr = 2, TT = 145 ns. Pozostałe parametry tych modeli mają wartości domyślne. Na podstawie literatury [27] przyjęto typowe wartości współczynników temperaturowych: $\alpha_{TT} = \alpha_{TD} = 3 \cdot 10^{-3} \text{ K}^{-1}$, $\alpha_{UD} = -2 \text{ mV/K}$. Rezystancje termiczne diody i tranzystora równe są 20 K/W. Częstotliwość kluczowania wynosi 100 kHz ($T_S = 10 \mu\text{s}$).

W metodzie z pracy [10] przyjęto wartości $R_{ON} = 0,6767\Omega$, $R_D = 0,12\Omega$, $R_D = 1\text{ V}$, wynikające z podanych powyżej wartości parametrów modeli tranzystora MOS i diody.

Na rys. 4 przedstawiono obliczone zależności napięcia wyjściowego i sprawności przetwornicy buck od współczynnika wypełnienia dla rezystancji obciążenia równej $R_0 = 3\Omega$. Na rys. 4–9 linie ciągłe oznaczają wyniki obliczeń uzyskane za pomocą zaproponowanej metody obliczeniowej, punkty — wyniki uzyskane z klasycznej elektrotermicznej analizy stanów przejściowych [17, 26], natomiast linie kreskowe — wyniki analiz izotermicznych uzyskanych metodą modeli uśrednionych z pracy [10].

Jak widać, uzyskano dobrą zgodność wyników analiz uzyskanych metodą analizy stanów przejściowych i metodą proponowaną w pracy. Wpływ samonagrzewania jest wyraźnie widoczny w zakresie dużych wartości współczynnika wypełnienia d , a różnice wartości napięcia wyjściowego uzyskanymi z analiz izotermicznych i elektrotermicznych dochodzą do 30%. Uzyskane przy użyciu modelu literaturowego [10] ujemne wartości sprawności w zakresie małych wartości współczynnika d są нефизyczne i wynikają z ograniczenia zakresu słuszności tego modelu do pracy przetwornicy w trybie ciągłym.

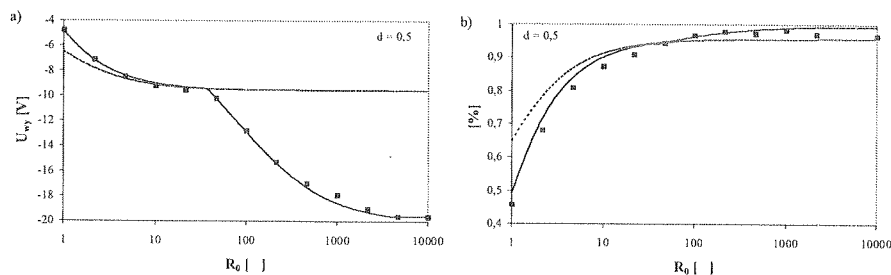


Rys. 4. Obliczone zależności napięcia wyjściowego (a) oraz sprawności (b) przetwornicy buck od współczynnika wypełnienia impulsów

Fig. 4. The calculated dependences of the output voltage (a) and the efficiency (b) of the buck converter on the duty factor

Z kolei, na rys. 5 przedstawiono zależność napięcia wyjściowego oraz sprawności energetycznej przetwornicy buck od rezystancji obciążenia przy współczynniku wypełnienia impulsów sterujących $d = 0,5$.

Przy rezystancji obciążenia $R_0 > 50\Omega$ układ pracuje w trybie nieciągłym, co skutkuje znaczącym wzrostem modułu napięcia wyjściowego i wzrostem sprawności energetycznej. W zakresie pracy nieciągłej wyniki analiz nieizotermicznych uzyskanych nową metodą i metodą analizy stanów przejściowych tylko nieznacznie różnią się między sobą, natomiast wyniki uzyskane przy wykorzystaniu modelu z pracy [10] różnią się od nich nawet o ponad 50% w przypadku napięcia wyjściowego i o ponad 30% w przypadku sprawności.



Rys. 5. Obliczone zależności napięcia wyjściowego (a) oraz sprawności (b) przetwornicy buck od rezystancji obciążenia

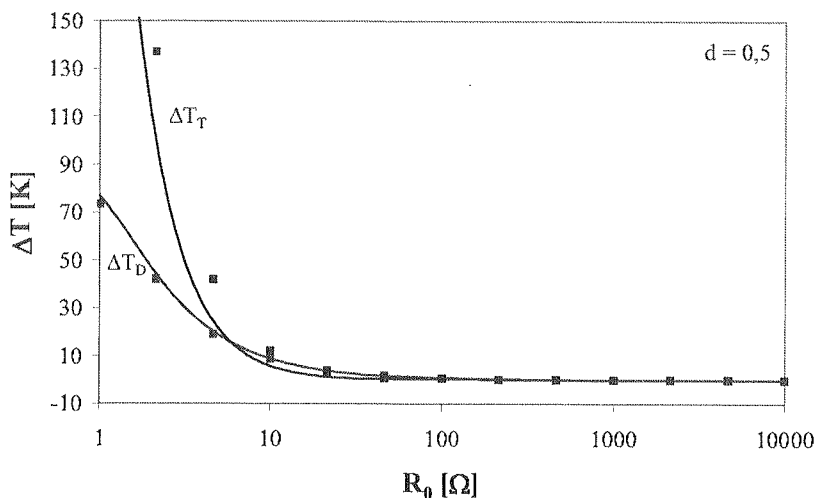
Fig. 5. The calculated dependences of the output voltage (a) and the efficiency (b) of the buck converter on the load resistance

Obliczone przy wykorzystaniu nowej metody i analizy stanów przejściowych zależności przyrostu temperatur wnętrza diody i tranzystora pokazano na rys. 6. Jak widać uzyskano dobrą zgodność wyników obliczeń obiema metodami dla diody i trochę gorszą dla tranzystora. Duże wartości przyrostów temperatury wnętrza są obserwowane dla małych rezystancji obciążenia.

Na rys. 7 pokazano obliczone zależności napięcia wyjściowego i sprawności przetwornicy boost od współczynnika wypełnienia dla rezystancji obciążenia równej $R_0 = 10 \Omega$. Widoczne są wyraźne różnice między charakterystykami izotermicznymi i nieizotermicznymi już dla współczynników wypełnienia $d > 0,5$. Różnice te przekraczają nawet 100% przy $d \approx 0,75$. Warto zauważyć, że dla $d > 0,8$, na skutek dużego spadku napięcia na włączonym kluczu tranzystorowym, przetwornica przestaje podwyższać napięcie. W tym zakresie napięcie wyjściowe stanowi różnicę napięcia wejściowego oraz spadków napięcia na diodzie i na rezystancji cewki. Zjawisko to nie jest uwzględnione w modelu literaturowym [10]. Sprawność przetwornicy maleje od blisko 90% do zaledwie 18%. Samonagrzewanie powoduje obniżenie sprawności przetwornicy.

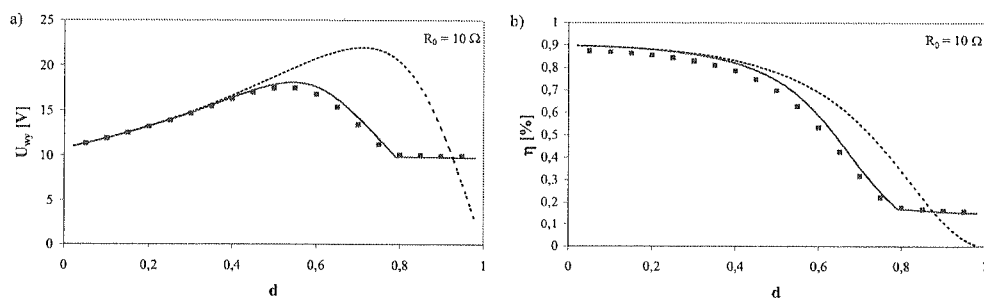
Z kolei, na rys. 8 przedstawiono zależność napięcia wyjściowego oraz sprawności energetycznej przetwornicy boost od rezystancji obciążenia przy współczynniku wypełnienia impulsów $d = 0,5$.

Przy rezystancji obciążenia $R_0 > 10 \Omega$ układ pracuje w trybie nieciągłym, co skutkuje znaczącym wzrostem napięcia wyjściowego i pogorszeniem sprawności energetycznej. W zakresie pracy nieciągłej wyniki analiz uzyskane nową metodą praktycznie pokrywają się z wynikami analizy stanów przejściowych, natomiast wyniki uzyskane przy wykorzystaniu modelu z pracy [10] różnią się od nich nawet o ponad 80% w przypadku napięcia wyjściowego i o ponad 30% w przypadku sprawności. Widoczne są różnice między wynikami uzyskanymi nową metodą i metodą stanów przejściowych w zakresie pracy nieciągłej, które rosną wraz ze wzrostem rezystancji R_0 , dochodząc do 10% na charakterystykach $U_{wy}(R_0)$ oraz do 5% na charakterystykach $\eta(R_0)$.



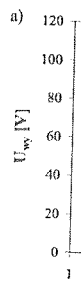
Rys. 6. Obliczone zależności przyrostu temperatur wnętrza tranzystora i diody ponad temperaturę otoczenia od rezystancji obciążenia

Fig. 6. The calculated dependences of the inner temperature excess of the transistor and the diode on the load resistance



Rys. 7. Obliczone zależności napięcia wyjściowego (a) oraz sprawności (b) przetwornicy boost od współczynnika wypełnienia impulsów

Fig. 7. The calculated dependences of the output voltage (a) and the efficiency (b) of the boost converter on the duty factor



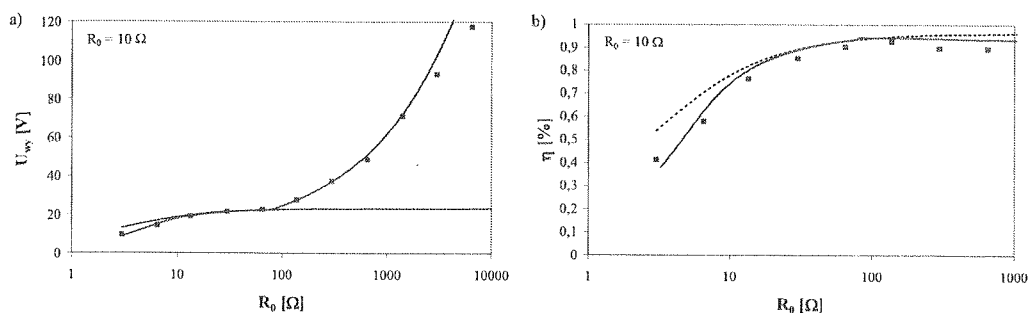
Rys.

Fig. 8. T

Rys. 9.

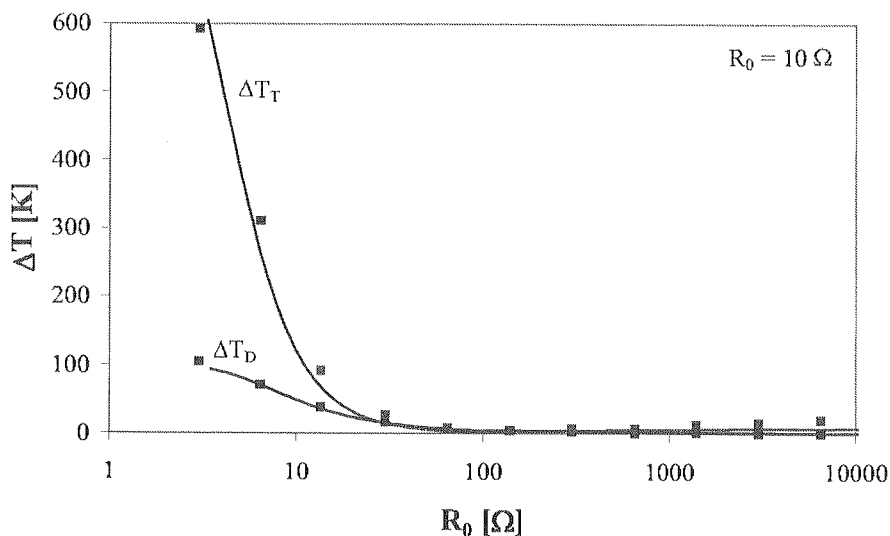
Fig. 9. The

Obli
ności prz
uzyskano
dla tranz
małych r



Rys. 8. Obliczone zależności napięcia wyjściowego (a) oraz sprawności (b) przetwornicy boost od rezystancji obciążenia

Fig. 8. The calculated dependences of the output voltage (a) and the efficiency (b) of the boost converter on the load resistance



Rys. 9. Obliczone zależności przyrostu temperatur wnętrza tranzystora i diody ponad temperaturę otoczenia od rezystancji obciążenia

Fig. 9. The calculated dependences of the inner temperature excess of the transistor and the diode on the load resistance

Obliczone przy wykorzystaniu nowej metody i analizy stanów przejściowych zależności przyrostu temperatur wnętrza diody i tranzystora pokazano na rys. 9. Jak widać uzyskano dobrą zgodność wyników obliczeń obiema metodami zarówno dla diody, jak dla tranzystora. Duże wartości przyrostów temperatury wnętrza są obserwowane dla małych rezystancji obciążenia.

Czas trwania obliczeń charakterystyk przy wykorzystaniu nowej metody wynosi około 0,1 s, podczas gdy wyznaczenie tych samych charakterystyk metodą analizy stanów przejściowych wynosi nawet kilka godzin.

5. PODSUMOWANIE

W pracy zaproponowano nowy sposób wyznaczania nieizotermicznych charakterystyk przetwornic dc-dc, bazujący na koncepcji uśrednionych modeli przetwornic dc-dc. Zaprezentowano ogólny sposób formułowania elektrotermicznego uśrednionego modelu przetwornic dc-dc o sterowaniu PWM oraz przedstawiono postać tego modelu dla przetwornic buck oraz boost.

W przeciwieństwie do literaturowych uśrednionych modeli przetwornic dc-dc, w zaproponowanej w niniejszej pracy metodzie uwzględniono wpływ nieliniowości charakterystyk elementów półprzewodnikowych, poziomów impulsów sterujących tranzystor kluczujący oraz samonagrzewania na charakterystyki rozważanych przetwornic w stanie ustalonym.

Przedstawione w poprzednim rozdziale wyniki obliczeń potwierdzają poprawność opracowanej metody, zarówno przy pracy rozważanych przetwornic w trybie ciągłego przepływu prądu, jak również w trybie nieciągłego przepływu prądu.

Zaletą opracowanej metody analizy jest możliwość określenia temperatur wnętrza elementów półprzewodnikowych, wyznaczenie napięcia wyjściowego i sprawności przetwornicy w obecności samonagrzewania. W porównaniu z metodą analizy stanów przejściowych uzyskano znaczące skrócenie czasu trwania obliczeń, który dla nowej metody nie przekracza 0,1 s, natomiast dla analizy stanów przejściowych przekracza kilka godzin. Przewaga nowej metody jest szczególnie wyraźna w zakresie dużych wartości rezystancji R_0 .

Wadą nowej metody jest pominięcie strat energii wydzielanej w elementach półprzewodnikowych w czasie przełączania, co skutkuje niewielkim zaniżeniem wartości ich temperatur wnętrza. Ogranicza to od góry zakres częstotliwości sygnału sterującego, przy których nowa metoda zapewnia dobrą dokładność obliczeń. Należy oczekiwać, że w sytuacji, gdy czasy przełączania diody i tranzystora staną się współmierne z okresem sygnału sterującego, nowa metoda nie zapewni żądanej dokładności obliczeń. W przypadku obecnie dostępnych tranzystorów mocy MOS oraz diod Schottky'ego, zakres słuszności metody będzie obejmował częstotliwości sygnału sterującego do około 1 MHz, natomiast dla tranzystorów bipolarnych — zakres ten będzie ograniczony do kilkudziesięciu kHz.

Przedstawiona metoda analizy przetwornic dc-dc jest uniwersalna, to znaczy jest ona słuszna zarówno dla pracy rozważanych układów w trybie ciągłego przepływu prądu, jak i w trybie przepływu nieciągłego oraz dla dowolnych rodzajów półprzewodnikowych elementów kluczujących.

1. C. no.
2. N. latia vol.
3. C. con
4. D. Inte
- pp.
5. A. 1992
6. C. Mot
7. S. I gnal Elec
8. D. A APE 2001
9. S. I Indus
10. R. E Acad
11. V.J. wer-f
12. Ch.P.
13. A. B
14. N. M and L
15. U. T
16. K. G gram
17. J. Z półprz Gdyni
18. J. Z a przetw

6. PODZIĘKOWANIA

Autor serdecznie dziękuje panu prof. Januszowi Zarębskiemu za cenne dyskusje i uwagi, które wpłynęły na poprawę czytelności rozważań prowadzonych w niniejszej pracy.

7. BIBLIOGRAFIA

1. C. Basso: *SPICE models verify SMPS designs*, Powerconversion & Intelligent Motion, vol. 27, no. 7, July 2001, pp. 26–30.
2. N. Mohan, W. P. Robbins, T. M. Undeland, R. Nilssen, O. Mo: *Simulation of Power Electronic and Motion Control Systems — An Overview*, Proceedings of the IEEE, vol. 82, Aug. 1994, pp. 1287–1302.
3. C. Basso: *Write your own generic SPICE power supplies controller models*, I. Guidelines. Powerconversion & Intelligent Motion, vol. 23, no. 4, April 1997, pp. 57–8, 60–2.
4. D. Kezic: *Simulation of DC-DC switch mode converters by means of personal computer*, 9th International Conference Electrical Drives and Power Electronics. KoREMA. 1996 Zagreb, Croatia, pp. 292–5.
5. A. Brambilla, E. Dallago: *A method for the simulation of switching DC-DC power supplies*, 1993 IEEE International Symposium on Circuits and Systems. Chicago, 1993 vol. 4, pp. 2664–7.
6. C. Basso: *Average simulations of flyback converters with SPICE3*, Powerconversion & Intelligent Motion, vol. 22, no. 12, Dec. 1996, pp. 10, 12–21.
7. S. Lineykin, S. Ben-Yaakov: *A unified SPICE compatible model for large and small signal envelope simulation of linear circuits excited by modulated signals*, IEEE 34th Annual Power Electronics Specialists Conference. Conference (PESC), Acapulco, 2003, vol. 3, pp. 1205–9.
8. D. Adar, S. Ben-Yaakov: *Generic average modeling and simulation of discrete controllers*, APEC 2001. Sixteenth Annual IEEE Applied Power Electronics Conference and Exposition, vol. 1, 2001, Anaheim, pp. 535–41.
9. S. Ben-Yaakov, I. Zehser: *PWM converters with resistive input*, IEEE Transactions on Industrial Electronics, vol. 45, no. 3, June 1998, pp. 519–20.
10. R. Ericson, D. Maksimovic: *Fundamentals of Power Electronics*, Norwell (USA), Kluwer Academic Publisher, 2001.
11. V.J. Thottuvelil, D. Chin, G.C. Verghese: *Hierarchical approaches to modelling high-power-factor ac-dc converters*, IEEE Trans. Power Electron, vol. 6, Apr. 1991, pp. 179–187.
12. Ch.P. Basso: *Switch-Mode Power Supply SPICE Cookbook*, New York, McGraw-Hill, 2001.
13. A. Borkowski: *Zasilanie urządzeń elektronicznych*, Warszawa, WKiŁ, 1990.
14. N. Mohan, T.M. Undeland, W.P. Robbins: *Power Electronics: Converters, Applications, and Design*, New York, John Wiley&Sons, 1995.
15. U. Tietze, Ch. Schenk: *Układy półprzewodnikowe*, Warszawa, WNT, 1987.
16. K. Górecki: *Przyspieszona analiza dławikowych przetwornic dc-dc o sterowaniu PWM w programie SPICE*, Kwartalnik Elektroniki i Telekomunikacji, Nr 4, 2005, s. 587.
17. J. Zarębski: *Modelowanie, symulacja i pomiary przebiegów elektrotermicznych w elementach półprzewodnikowych i układach elektronicznych*, Prace Naukowe Wyższej Szkoły Morskiej w Gdyni, Gdynia, 1996.
18. J. Zarębski, K. Górecki: *Analiza wpływu samonagrzewania na charakterystyki dławikowych przetwornic dc-dc*, Elektronika, Not-Sigma, Warszawa, Nr 2–3, 2003, s. 24.

19. K. Górecki, J. Zarębski: *Modelowanie wpływu efektów termicznych na właściwości przetwornicy BOOST w programie SPICE*, VI Krajowa Konferencja Naukowa Sterowanie w Energoelektronice i Napędzie Elektrycznym SENE 2003, Łódź, 2003, t. 1, s. 143.
20. K. Górecki, J. Zarębski: *Wyznaczanie charakterystyk przetwornicy BOOST w stanie ustalonym*, X Konferencja Naukowo-Techniczna Zastosowania Komputerów w Elektrotechnice ZKwE'2005, Poznań, 2005, s. 251.
21. D. Stanny: *Badanie wpływu parametrów elementów kluczujących na charakterystyki wybranych przetwornic dc-dc*, Praca dyplomowa magisterska, Wydział Elektryczny, Akademia Morska w Gdyni, 2004.
22. K. Górecki, J. Zarębski: *Wyznaczanie temperatury wnętrza półprzewodnikowych elementów mocy sterowanych sygnałem w.cz.*, XXVIII Międzynarodowa Konferencja z Podstaw Elektrotechniki i Teorii Obwodów IC-SPETO'2005, Ustroń, 2005, Vol. 2, s. 253.
23. K. Zamłyński: *Modelowanie termiczne sprzężenia zwrotnego w monolitycznych układach scalonych przy pomocy programu SPICE-2*, Rozprawy Elektrotechniczne, 1990, 35, z. 3, s. 673.
24. K. Górecki: *Elektrotermiczny makromodel tranzystora Darlingtona do analizy układów elektrotermicznych*, Praca doktorska, Politechnika Gdańska, 1999.
25. K. Górecki, J. Zarębski: *The SPICE Macromodel of MC34066 Controller*, 6th International Conference on Unconventional Electromechanical and Electrical Systems UEES'04, Alushta, 2004, Vol. 2, p. 525.
26. K. Górecki, J. Zarębski, K. Posobkiewicz: *The Electrothermal Macromodel of the Monolithic Voltage Regulator LT1073 for SPICE*, 10th IEEE International Conference on Electronics, Circuits and Systems ICECS 2003, Sharjah, 2003, p. 1113.
27. W. Kuźmich: *Projektowanie analogowych układów scalonych*, WNT, Warszawa, 1985.

K. GÓRECKI

CALCULATIONS OF THE NONISOTHERMAL CHARACTERISTICS OF DC-DC CHOPPERS WITH PWM CONTROL AT STEADY-STATE USING AVERAGE MODELS METHOD

S u m m a r y

The dc-dc choppers are the base components of switched voltage regulators. For the analysis of such a class of electronic circuits the classical transient analysis or the dc analysis with the use of averaged models of such converters is applied. The average models method makes it possible to obtain the characteristics of the considered converters at steady state in less calculation time than in the transient analysis method.

The literature averaged models of dc-dc converters were elaborated by neglecting the selfheating phenomenon and nonlinear characteristics of semiconductor devices.

In this paper a new method of calculating the dc-dc converters characteristics with PWM control in steady state with selfheating in semiconductor devices taken into account, is proposed. This method is based on the conception of the averaged models method and it enables the calculations of the values of such dc-dc converters parameters as the output voltage and the efficiency of the converter, as well as the inner temperatures of the diode and the transistor, which are important from the point of view of the reliability of the converters. In the elaborated method, realized with the use of SPICE, the averaged electrothermal model of dc-dc converter is formulated. This model consists of three kinds of circuits. The first one constitutes the electrical network composed of the controlled voltage sources and resistors. The structure of this circuit results from the principles of formulating the averaged models of dc-dc converters. The second kind of circuits includes the series connection of the transistor and the diode described with the use of built-in SPICE models of the considered devices and the controlled voltage sources, which

model the temperature influence on the voltage across the switched-on semiconductor devices. The last circuits are the electrical analogues of the thermal models of the transistor and the diode.

The method formulating the considered circuits on the example of buck and boost converters is presented. The correctness of the elaborated method was verified by comparing the calculation results obtained from the new method and from the classical transient analysis. The good agreement between the considered methods is obtained, which proves the correctness of the proposed method.

Keywords: dc-dc converters, electrothermal analysis, SPICE, average models

Para

me
tw
me

Sta

Wsp
udostęp
ków logi
do dzies
gigabitov
Pro i Vir
Przykład
jących w
W w
powszech
ku jak i
[13-16].
Xilinx) u

Parametryzowany interfejs komunikacyjny dla układów FPGA

KRZYSZTOF POŹNIAK

*Instytut Systemów Elektronicznych,
Politechnika Warszawska,
ul. Nowowiejska 15/19, 00-665 Warszawa,
e-mail: Pozniak@ise.pw.edu.pl.*

*Otrzymano 2005.12.03
Autoryzowano 2005.12.03*

W pracy opisano warstwę sprzętową parametryzowanego, uniwersalnego interfejsu komunikacyjnego do zastosowań w układach FPGA. Omówiono metodykę automatycznego tworzenia przestrzeni adresowej i danych zgodnie z deklaracjami użytkownika. Opisano metody standaryzacji komunikacji I/O oraz przedstawiono przykłady implementacji.

Słowa kluczowe: Altera, Xilinx, VHDL, parametryzacja, interfejsy komunikacyjne, programowanie behawioralne

1. WSTĘP

Współczesna technologia Field-Programmable Gate Array (FPGA) [1-5], może udostępniać w jednym układzie scalonym do kilkuset tysięcy konfigurowalnych bloków logicznych, ponad dwieście szybkich bloków przetwarzania numerycznego [6-9], do dziesięciu megabajtów pamięci statycznej, kilkanaście niezależnych modułów dla gigabitowej transmisji łączami elektrycznymi lub optycznymi [10-12]. Serie VirtexII-Pro i Virtex4 firmy Xilinx zawierają dodatkowo wbudowane procesory klasy Power-PC. Przykłady dostępnych na rynku generacji najnowszych serii układów FPGA, z przodujących w tej dziedzinie firm Altera i Xilinx, przedstawiono w tabeli 1.

W wyniku ciągłego rozwoju możliwości funkcjonalnych, układy FPGA są coraz powszechniej stosowane zarówno w popularnych urządzeniach powszechnego użytku jak i w rozległych, wielokanałowych specjalizowanych systemach elektronicznych [13-16]. Oferowane narzędzia projektowe (np. Quartus firmy Altera lub ISE firmy Xilinx) udostępniają konstruktorom zintegrowane środowisko umożliwiające realizację

Tabela 1

Porównanie parametrów porównywalnych serii układów FPGA firm Altera i Xilinx

Comparison of parameters for similar families of FPGA chips by Altera and Xilinx

Firma/seria	Bloki logiczne	Pamięć RAM	Mnożniki 18x18 bitów	Transmisja szeregową	Nóżki użytkownika
Xilinx/Virtex4	142128	9936 kB	192	24 (3.125Gb/s)	896
Altera/Stratix IIGX	132,540	6747 kB	252	20 (6.375 Gb/s)	734

zaawansowanych projektów wieloukładowych (w językach VHDL, Verilog, C itp.), symulację funkcjonalną i pokompilacyjną, analizę czasową i obciążeniową oraz wiele innych. Technologia FPGA wypiera układy ASIC dzięki upowszechnieniu oferty rynkowej dla produkcji seryjnych rozwiązań projektowych typu „HardCopy” [2].

Wykorzystywanie układów FPGA w urządzeniach elektronicznych powoduje, że można w nich łatwo i szybko realizować modyfikacje w warstwie funkcjonalnej, bez konieczności dokonywania jakichkolwiek zmian w istniejącej strukturze sprzętowej lub realizacji nowej wersji urządzenia. W związku z tym, urządzenia z układami FPGA, są bardzo często wyposażone w rozbudowane interfejsy komunikacyjne. Umożliwiają one w pełni zdalne zarządzanie urządzeniem, jego konfigurację, szczegółową diagnostykę i monitoring [14]. Zmiany w warstwie funkcjonalnej muszą być poprawnie odwzorowane w warstwie komunikacyjnej zarówno po stronie sprzętu w układach FPGA, jak i oprogramowania zarządzającego.

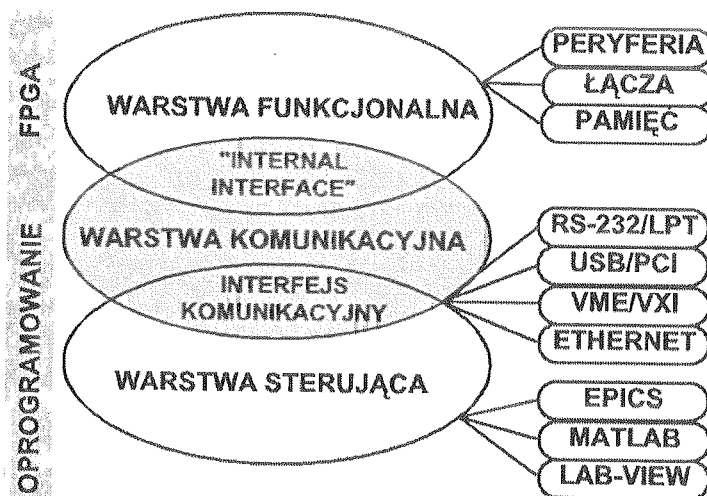
Profesjonalne, zintegrowane systemy komunikacji I/O pomiędzy środowiskiem programistycznym a układami FPGA oferuje wiele komercyjnych firmy, jak np.:

- National Instrument, który opracował technologię integracji „LabView” do bezpośredniej współpracy z National Instruments Reconfigurable I/O (RIO) devices [17],
- MathWorks, oferujący narzędzia MATLAB i SIMULINK bardzo dobrze dostosowane do współpracy z wieloma urządzeniami komercyjnymi [18],
- Nallatech, producent zaawansowanych urządzeń zawierających układy FPGA i DSP najnowszych generacji, opracował środowisko programistyczne FUSE [19],
- Xilinx, wytwórca układów FPGA, oferuje narzędzia programistyczne współpracujące z urządzeniami elektronicznymi wyposażonymi w FPGA od innych wytwórców [1],
- Altera, wytwórca układów FPGA, oferuje narzędzia programistyczne SOPC Builder współpracujące z magistralą „Avalon” [20].

Należy w tym miejscu również wymienić uniwersalną magistralę „Wishbone”, z którą współpracuje wiele urządzeń opracowanych na licencji „Open Core” [21].

W niniejszej pracy przedstawiono zrealizowane w praktyce przez autora, spójne, systemowe rozwiązanie interfejsu komunikacyjnego „Internal Interface” [22, 23], nazywanego dalej w pracy Interfejsem. Interfejs jest automatycznie tworzony i implementowany w układach FPGA i może być sprzężony z warstwą programistyczną komputerowego systemu sterowania. Ogólny schematstruktury warstwowej urządzeń z interfejsem komunikacyjnym implementowanym w FPGA przedstawiono na rys 1. Interfejs stanowi pomost pomiędzy warstwą funkcjonalną układu FPGA z układami

zewnętrznymi, a fizyczną warstwą komunikacyjną zarządzaną przez programistyczną warstwę sterującą.



Rys. 1. Warstwowa struktura urządzenia FPGA z interfejsem komunikacyjnym

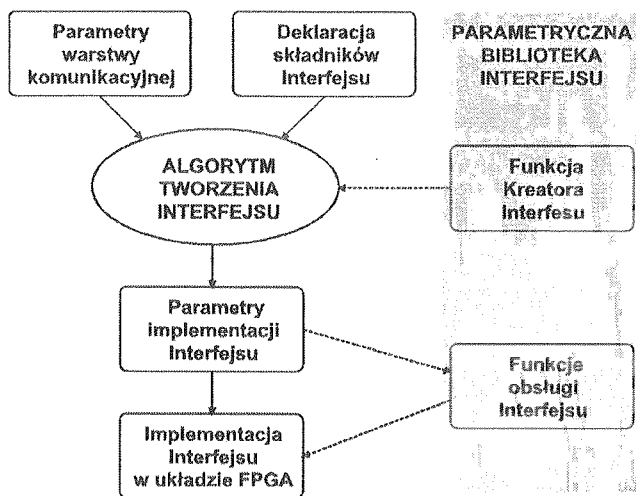
Fig. 1. Layered structure of FPGA device with communication interface

2. STRUKTURA FUNKCJONALNA INTERFEJSU DLA UKŁADU FPGA

Struktura funkcjonalna Interfejsu ma za zadanie umożliwić automatyzację projektowania lokalnego interfejsu komunikacyjnego w układzie FPGA oraz jego sprzężenie z fizyczną warstwą komunikacyjną. Ogólną zasadę procesu projektowania Interfejsu oraz powiązań funkcjonalnych w implementacji pokazano na rys. 2.

Opis struktury implementowanego Interfejsu jest umieszczony w strukturze definicyjnej, w formie deklaracji *składników Interfejsu*. Proces tworzenia fizycznej implementacji jest przeprowadzany automatycznie za pomocą *algorytmu tworzenia Interfejsu*. Istotną rolę w tym procesie odgrywają parametry warstwy komunikacyjnej, które są formalnie niezależne od deklaracji składników Interfejsu. Powyższe rozwiązanie funkcjonalne umożliwia opracowanie projektu z Interfejsem w sposób niezależny od fizycznej warstwy komunikacyjnej. Takie rozwiązanie pozwala na realizację uniwersalnych bibliotek komponentów współpracujących z Interfejsem, które są przenaszalne i współpracują z różnymi standardami warstwy komunikacyjnej.

Warstwa sprzętowa została zrealizowana w sposób jednolity, w formie *parametrycznej biblioteki Interfejsu* opracowanej w języku VHDL. Funkcja *Kreatora Interfejsu* wypracowuje rzeczywiste parametry Interfejsu. Na ich podstawie, biblioteczny zestaw *funkcji obsługi Interfejsu*, umożliwia jego automatyczną implementację oraz



Rys. 2. Schemat tworzenia i implementacji Interfejsu w układzie FPGA

Fig. 2. Creation and implementation method of the Interface in FPGA circuit

powiązania z elementami zewnętrznymi. W następnych rozdziałach opisano metodykę projektowania oraz przykłady implementowania Interfejsu.

2.1. DEKLARACJA SKŁADNIKÓW INTERFEJSU

Interfejs jest deklarowany w formie listy rekordów. Pojedynczy rekord listy stanowi uporządkowane składniki. Parametry składnika dzielą się na:

- identyfikacyjne, umożliwiające jednoznaczne rozróżnienie rekordu (typ, nazwa),
- skalujące, definiujące fizyczne rozmiary rekordu,
- dostęp, określające prawa dostępu do rekordu w trybie zapisu lub odczytu

Poprzez dobór składników i nadanie ich parametrom określonych wartości, użytkownik może kształtować strukturę Interfejsu zgodnie z wymogami danego projektu. Najważniejsze właściwości kształtujące Interfejs zostały wymienione poniżej:

1. Identyfikator składnika należy do grupy parametrów identyfikujących. Stanowi unikalny identyfikator rekordu. W celu uzyskania przejrzystości projektu, jako identyfikatory są wykorzystywane odrębne stałe symboliczne (mnemoniki), definiowane przez użytkownika. Ich użycie w sposób jednoznaczny wskazuje na odniesienie do przypisanego mu składnika.
2. Typ składnika należy do grupy parametrów identyfikujących. Określa typ pojedynczego rekordu poprzez przypisanie go do jednej z dwóch grup rodzajowych:
 - *fizycznych*, definiujących rzeczywiste obiekty Interfejsu jak: obszary pamięci, rejestry lub pojedyncze bity. Przyjęte rozwiązanie obejmuje powszechnie wykorzystywane struktury funkcjonalne implementowane w układach FPGA,

- *grupujących*, umożliwiających tworzenie wspólnych obszarów adresowania oraz łączenia bitów we wspólny wektor. Są to składniki wirtualne, umożliwiające użytkownikowi ustalenie własnej organizacji Interfejsu w określonej jego części lub w całości. Grupowanie odbywa się poprzez powiązanie identyfikatorów składników fizycznych z danym rekordem grupującym.
- 3. Wymiary składnika fizycznego należą do grupy parametrów skalujących. *Parametr szerokości elementu* określa liczbę bitów pojedynczego elementu składnika, natomiast *parametr liczby elementów* wyznacza rozmiar tablicy złożonej z identycznych elementów. Dwuwymiarowe skalowanie w sposób jednolity obejmuje wszystkie składniki fizyczne. Umożliwia realizację projektów parametrycznych, gdzie rozmiary i liczba elementów Interfejsu jest definiowana symbolicznie na etapie deklaracji.
- 4. Prawa dostępu do rekordu określają zarówno prawo zapisu jak i prawo odczytu zawartości danego rekordu poprzez magistralę Interfejsu. Definiują także, czy proces dostępu odbywa się do niezależnego elementu peryferyjnego dołączonego do Interfejsu, czy do elementu rejestrowego implementowanego wewnątrz Interfejsu.

2.2. ALGORYTM TWORZENIA INTERFEJSU

Tworzenie fizycznej implementacji Interfejsu w układzie FPGA odbywa się automatycznie. Proces ten jest przeprowadzany przez funkcję biblioteczną *Kreatora Interfejsu* z wykorzystaniem sparametryzowanego algorytmu analizy opisanego w języku VHDL. Jako argumenty funkcji są wprowadzane: lista deklaracji składników (por. rozdz. 2.1) oraz parametry warstwy komunikacyjnej.

W niniejszym rozdziale przedstawiono zasady tworzenia Interfejsu w zakresie grupowania, dopasowania do parametrów magistrali komunikacyjnej, wypełnienia przestrzeni adresowej, podziału rekordów fizycznych na części w przypadku przekraczania szerokości magistrali danych Interfejsu itp. Finalnym efektem tego procesu jest automatyczne utworzenie fizycznej implementacji Interfejsu, tzn. *wypracowanie parametrów Interfejsu* (por. rys. 2).

Parametry magistrali komunikacyjnej, uwzględniane w procesie tworzenia Interfejsu, definiują:

- A — liczbę linii magistrali adresowej wyznaczającą dostępną ciągłą przestrzeń adresową w zakresie pozycji adresowych liczonych od wartości 0 do $2^A - 1$,
- D — szerokość słowa danych wyrażoną w bitach. Wartości przesyłanych danych zawiera się w przedziale od 0 do $2^D - 1$.

Należy podkreślić, że proces konwersji niezależnia wirtualną przestrzeń zajmowaną przez składniki Interfejsu od rzeczywistej przestrzeni adresowej i danych udostępnianej w warstwie komunikacyjnej. W ten sposób konstrukcja projektu jest niezależniona od fizycznej budowy magistrali komunikacyjnej, a deklarowane składniki Interfejsu odzwierciedlają wyłącznie funkcjonalne wymagania projektu.

Rozmieszczenie w fizycznej przestrzeni adresowej tablic wektorowych (np. tablicy rejestrów implementowanych w blokach LCELL) polega na ich konwersji z przestrze-

ni wirtualnej, zadanej przez parametry skalujące. W tab. 2 zamieszczono przykład konwersji tablicy trzech rejestrów 18-bitowych (oznaczonych jako R0, R1, R2) na przestrzeń 8-mio bitowej magistrali danych Interfejsu. Przyjęto, że adresowanie jest inicjalizowane od pozycji 0.

Tabela 2

Przykład dekompozycji tablicy rejestrów w przestrzeni adresowej magistrali komunikacyjnej
(szare pola oznaczają niewykorzystane bity danych)

Example of decomposition of register table in the address space of communication bus
(grey fields mean non used data bits)

Adres	D7	D6	D5	D4	D3	D2	D1	D0	uwagi
0	R0-7	R0-6	R0-5	R0-4	R0-3	R0-2	R0-1	R0-0	rejestr R0
1	R0-15	R0-14	R0-13	R0-12	R0-11	R0-10	R0-9	R0-8	
2							R0-17	R0-16	
3	R1-7	R1-6	R1-5	R1-4	R1-3	R1-2	R1-1	R1-0	rejestr R1
4	R1-15	R1-14	R1-13	R1-12	R1-11	R1-10	R1-9	R1-8	
5							R1-17	R1-16	
6	R2-7	R2-6	R2-5	R2-4	R2-3	R2-2	R2-1	R2-0	rejestr R2
7	R2-15	R2-14	R2-13	R2-12	R2-11	R2-10	R2-9	R2-8	
8							R2-17	R2-16	

Podział 18-to bitowego rejestru na 8-mio bitowe części wymaga rezerwacji trzech kolejnych pozycji adresowych w przestrzeni Interfejsu. W konsekwencji, trzy kolejne rejestry tablicy zajmują łącznie dziewięć adresów.

Specjalną formę dekompozycji przestrzeni adresowej zastosowano dla obszarów pamięci. Są one realizowane w układach FPGA poprzez implementację kilku wzajemnie sprzężonych bloków pamięci SRAM. Liczba i rozmiary wykorzystanych bloków pamięci wynikają bezpośrednio z zadanych parametrów skalujących. W tab. 3 zamieszczono przykład konwersji trzech komórek pamięci o szerokości słowa 20-bitów na przestrzeń 8-bitowej magistrali danych Interfejsu. Przyjęto, że adresowanie jest inicjalizowane od pozycji 7.

Podział 20-to bitowego słowa pamięci na 8-mio bitowe części wymaga wykorzystania trzech odrębnych bloków pamięci SRAM. Zostały one oznaczone w tab. 2 jako: A, B i C. Każdy blok pamięci posiada zarezerwowaną własną przestrzeń adresową w przestrzeni Interfejsu w taki sposób, aby najmłodsze linie adresowe magistrali komunikacyjnej mogły być bezpośrednio podłączone do obszarów pamięci. Stąd dla powyższego przykładu dokonano rezerwacji dwóch najmłodszych linii adresowych (A0 i A1). W konsekwencji, ostatnia pozycja adresowania każdego bloku pamięci pozostaje niewykorzystana. Wyboru jednego z trzech bloków pamięci SRAM dokonują kolejne dwie starsze linie adresowe (A2 i A3), stąd ostatnia pozycja adresowania bloku pamięci pozostaje zarezerwowana, ale nie jest niewykorzystana (pozycje adresowe: 24-27). Uzyaskanie warunku początkowego adresowania (tzn. adresy A3-0=0) wymagało osadzenia pamięci dopiero od adresu 16.

Str

wania s
przez p
adresow
mentacj
uproszc
niej ws

Mag
cyjnej (p
• linie n
• linie n
• linie s
• II.R
Sygn
mięc
• II.O
• II.W
• II.ST
wewn
dany

zykład
R2) na
ie jest

Tabela 2

ej

0

1

2

trzech
kolejne

oszarów
wzajem-
bloków
. 3 za-
0-bitów
jest ini-

zysztania
ako: A,
sową w
i komu-
dla po-
ch (A0 i
ozostaje
kolejne
pamięci
7). Uzy-
sadzenia

Tabela 3

Przykład dekompozycji obszaru pamięci w przestrzeni adresowej magistrali komunikacyjnej
(szare pola oznaczają niewykorzystane bity danych)

Decomposition example of memory area for communication bus
(grey fields mean non used data bits)

Adres	A3	A2	A1	A0	D7	D6	D5	D4	D3	D2	D1	D0	uwagi
7-15													
16	0	0	0	0	A0-7	A0-6	A0-5	A0-4	A0-3	A0-2	A0-1	A0-0	Pamięć A
17	0	0	0	1	A1-7	A1-6	A1-5	A1-4	A1-3	A1-2	A1-1	A1-0	
18	0	0	1	0	A2-7	A2-6	A2-5	A2-4	A2-3	A2-2	A2-1	A2-0	
19	0	0	1	1									
20	0	1	0	0	B0-7	B0-6	B0-5	B0-4	B0-3	B0-2	B0-1	B0-0	Pamięć B
21	0	1	0	1	B1-7	B1-6	B1-5	B1-4	B1-3	B1-2	B1-1	B1-0	
22	0	1	1	0	B2-7	B2-6	B2-5	B2-4	B2-3	B2-2	B2-1	B2-0	
23	0	1	1	1									
24	1	0	0	0					C0-3	C0-2	C0-1	C0-0	Pamięć C
25	1	0	0	1					C1-3	C1-2	C1-1	C1-0	
26	1	0	1	0					C2-3	C2-2	C2-1	C2-0	
27	1	0	1	1									
28-31													

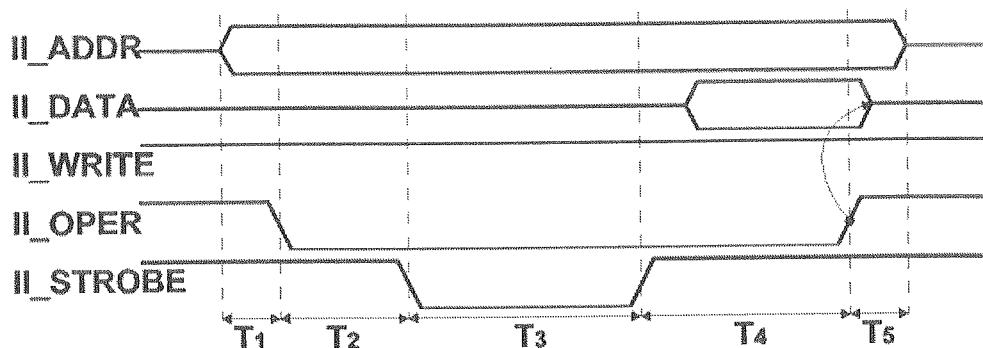
Stronicowanie przestrzeni adresowej Interfejsu jest powiązane z procesem grupowania składników. Każdy składnik przynależy do jednej z wcześniej zadeklarowanych przez projektanta grup. Każda grupa uzyskuje automatycznie wyodrębnioną przestrzeń adresową. Wyróżnikiem grupy jest jej własny prefiks adresowy, który w procesie implementacji w układzie FPGA przyspiesza proces dekodowania adresów. Dla rozwiązania uproszczonego, wymagana jest obligatoryjnie deklaracja jednej grupy i dołączenie do niej wszystkich zadeklarowanych składników.

2.3. PARAMETRIZOWANA MAGISTRALA KOMUNIKACYJNA INTERFEJSU

Magistrala komunikacyjna Interfejsu odzwierciedla parametry warstwy komunikacyjnej (parametry A i D — por. rozdz. 2.2). Jest podzielona na trzy grupy sygnałów:

- linie magistrali adresowej **IL_ADDR** o szerokości A bitów,
- linie magistrali danych, **IL_DATA** o szerokości D bitów.
- linie sterowania, pozwalające realizować operacje dostępu oraz inicjalizacji:
 - **IL_RESET** — stan niski wymusza asynchroniczny proces inicjalizacji Interfejsu. Sygnał ten jest również wykorzystywany do inicjalizowania rejestrów, bloków pamięci i innych bloków funkcjonalnych umieszczonych wewnątrz układu FPGA,
 - **IL_OPER** — stan niski oznacza wykonywanie operacji dostępu.
 - **IL_WRITE** — stan niski oznacza wykonywanie operacji zapisu, a wysoki odczytu,
 - **IL_STROBE** — zbocze opadające oznacza ważny adres na magistrali adresowej wewnątrz układu FPGA, a zbocze narastające oznacza ważne dane na magistrali danych w operacji zapisu.

Ogólną sekwencję czasową pojedynczej operacji odczytu danych z peryferyjnego układu FPGA zamieszczono na rys. 3.



Rys. 3. Sekwencja operacji odczytu z układu FPGA

Fig. 3. Read operation sequence from FPGA circuit

Poszczególne etapy sekwencji odczytu wyznaczają kolejne przedziały czasu $T_1 - T_5$:

- Na etapie T_1 ważny adres ustawiony przez sterownik jest dystrybuowany do bloków dekodowania układu FPGA.
- Na etapie T_2 układ FPGA otwiera kanał dystrybucji danych na podstawie ważnego adresu i niskiego stanu sygnału **II_OPER**. Etap T_2 kończy opadające zbocze sygnału **II_STROBE**, które służy do rejestracji ważnego adresu, m.in. w blokach pamięci synchronicznych umieszczonych w najnowszych seriach układów FPGA.
- Na etapie T_3 następuje wypracowanie ważnych danych. Etap T_3 kończy narastające zbocze sygnału **II_STROBE**, które służy do taktowania wyjścia danych w układach synchronicznych, m.in. w blokach pamięci synchronicznych.
- Etap T_4 jest przeznaczony na wyprowadzenie ważnych danych z układu FPGA i ich dystrybucję poprzez magistralę komunikacyjną do sterownika. Etap T_4 kończy zmianę stanu sygnału **II_OPER** na wysoki. Układ FPGA niezwłocznie przerywa proces wystawiania danych.
- Na zakończenie etapu T_5 sterownik zdejmuję ważny adres z magistrali adresowej.

Ogólną sekwencję czasową pojedynczej operacji zapisu danych do peryferyjnego układu FPGA zamieszczono na rys. 4.

Poszczególne etapy sekwencji zapisu wyznaczają, podobnie jak dla odczytu, kolejne przedziały czasu $T_1 - T_5$. Dane przesyłane ze sterownika są rejestrowane w układzie FPGA narastającym zboczem sygnału **II_STROBE** na zakończenie przedziału czasu T_3 . Rejestracji podlegają zarówno dane w prostych przerzutnikach jak i synchronicznych blokach pamięci.

W rozwiązaniach praktycznych o szybkości działania Interfejsu decydują nie tylko parametry czasowe magistrali komunikacyjnej, ale również uzyskana szybkość działania Interfejsu wewnątrz układu FPGA. W tab. 4 podano orientacyjne przedziały czasu dla poszczególnych etapów sekwencji operacji dostępu dla powszechnie sto-

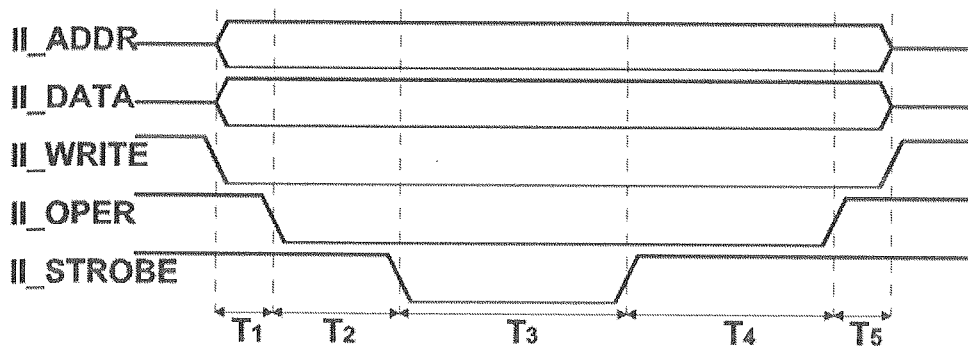
II_A
II_D
II_V
II_C
II_S

sowany
Xilinx
rzeczy

Cza
Czas

Op
asynch
mentow
nego w
przerzu
sygnał
rozwiąz
zminim

Ob
standar
cji do b
2) i odb
składni
na pozi



Rys. 4. Sekwencja operacji zapisu do układu FPGA

Fig. 4. Write operation sequence to FPGA circuit

sowanych obecnie serii układów firmy Altera (ACEX, CYCLONE, STRATIX) oraz Xilinx (SPARTAN, VIRTEX). Istotnym czynnikiem wprowadzającym opóźnienia jest rzeczywisty poziom złożoności Interfejsu wynikający z deklaracji jego składników.

Tabela 4

Orientacyjne przedziały czasu dla nowych serii układów FPGA firm Altera i Xilinx

Approximate time slots for new FPGA circuits by Altera and Xilinx

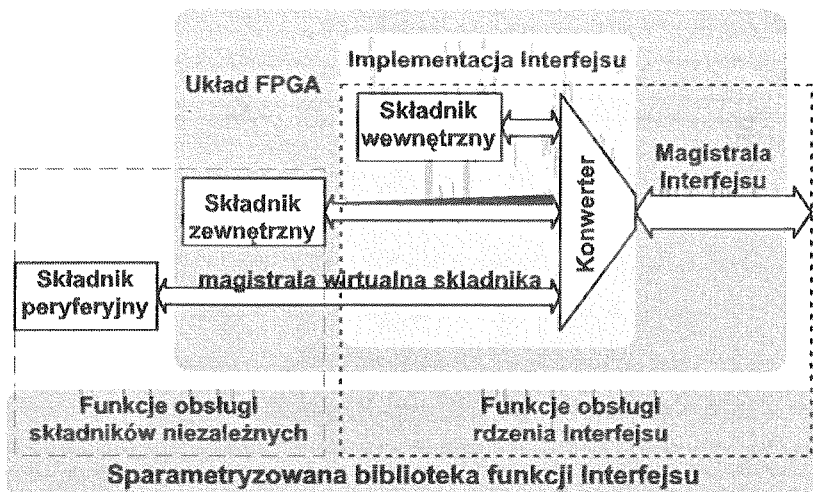
etap	T ₁	T ₂	T ₃	T ₄	T ₅	Σ
Czas minimalny [ns]	10	15	10	10	5	50
Czas maksymalny [ns]	25	40	50	50	10	175

Operacja dostępu poprzez Interfejs jest realizowana w układzie FPGA jako proces asynchroniczny i nie wymaga stosowania dodatkowego sygnału zegarowego. Zaimplementowano funkcję sygnału **II_STROBE** jako sygnału zegara globalnego dystrybuowanego wewnątrz układu FPGA. Jest on podawany na wszystkie elementy rejestrowe, jak przerzutniki i pamięci synchroniczne. Sygnał **II_RESET** jest implementowany jako sygnał globalny, asynchronicznie kasujący wszystkie elementy rejestrowe. Powyższe rozwiązanie w istotny sposób pozwala zwiększyć szybkość działania Interfejsu oraz zminimalizować liczbę wewnętrznych połączeń w układzie FPGA.

2.4. STANDARDYZOWANA OBSŁUGA IMPLEMENTACJI INTERFEJSU

Obsługa implementacji Interfejsu w układzie FPGA jest realizowana w sposób standaryzowany, poprzez wydzielone funkcje biblioteczne. Dostosowanie działania funkcji do bieżącej struktury Interfejsu wynika z jego aktualnych parametrów (por. rozdz. 2) i odbywa się automatycznie. Modyfikacja struktury Interfejsu na poziomie deklaracji składników nie wymaga modyfikacji projektu na poziomie obsługi Interfejsu, a jedynie na poziomie obsługi zmienionych funkcjonalności (np. realizacji procesu dla rejestru

dodanego w deklaracji składników). W konsekwencji powyższe rozwiązanie obejmuje dwa wydzielone etapy zadaniowe obsługi Interfejsu, co schematycznie pokazano na rys. 5.



Rys. 5. Schemat funkcjonalny implementacji Interfejsu w układzie FPGA z wykorzystaniem sparametryzowanych funkcji bibliotecznych

Fig. 5. Functional diagram of Interface implementation in FPGA circuit with usage of parameterized library functions

Zestaw funkcji obsługi rdzenia Interfejsu jest niezależny od implementowanej struktury Interfejsu i zapewnia:

- sterowanie magistralą komunikacyjną Interfejsu,
- zarządzanie blokiem Konwertera Magistrali Interfejsu na magistrale wirtualne dla poszczególnych składników,
- wykonanie procesów zapisu i odczytu dla składników implementowanych automatycznie wewnątrz Interfejsu.

Funkcje obsługi składników niezależnych obsługują składniki zewnętrzne i peryferyjne. Dobór odpowiednich funkcji obsługi wynika z typu składnika oraz zdefiniowanych właściwości dostępu do niego poprzez Interfejs. W bibliotece Interfejsu umieszczono zestawy funkcji do obsługi wszystkich typów deklarowanych składników fizycznych.

3. BADANIE IMPLEMENTACJI INTERFEJSU W UKŁADACH FPGA

W niniejszym rozdziale przedstawiono przykład prostej realizacji Interfejsu dla różnych parametrów magistrali komunikacyjnej. Implementacje Interfejsu wykonano dla wybranych układów FPGA. Wchodzą one w skład nowych serii firm Altera i

muje
no na

Xilinx. Uzyskane rezultaty przedstawiono w tabeli porównawczej. Badaniu poddano strukturę Interfejsu przedstawioną w tab. 5. Zadeklarowano:

- dwa rejestry: pojedynczy 8-bitowy **REJESTR1** i podwójny 10-bitowy **REJESTR2**,
 - 3 elementy bitowe (**BITY1**, **BITY2** i **BITY3**) zgrupowane jako **WEKTOR**. **BITY1** są składnikiem dwubitowym, **BITY2** są zadeklarowane jako jeden bit, a **BITY3** są dwuelementową tablicą, gdzie każdy element zawiera 2 bity danych,
 - 3 komórki obszaru pamięci (**PAMIĘĆ**) szerokości 6-bitów danych każda.
- Elementy rejestrowe zgrupowano jako **STRONA1**, a pamięć jako grupę **STRONA2**.

Tabela 5

Przykładowa struktura deklaracji Interfejsu

Exemplary structure of Interface declaration

Strona	Wektor	Mnemonik	Szerokość [bity]	Liczba	Uwagi
STRONA1		REJESTR1	8	1	Rejestr 8-bitowy
		REJESTR2	6	1	Rejestr 10-bitowy
	WEKTOR	BITY1	2	1	2 bity danych
	WEKTOR	BITY2	1	1	1 bit danych
	WEKTOR	BITY3	2	2	2 razy 2 bity danych
STRONA2		PAMIĘĆ	6	3	3 komórki 6-bitowe

ized

Obraz implementacji dla 4-bitowej magistrali Interfejsu zamieszczono w tab. 6. Dane składników zostały podzielone na części i umieszczone w obszarze adresowania.

struk-

Tabela 6

Implementacja przykładu dla 4-bitowej magistrali Interfejsu

Implementation of an example for 4-bit Interface bus

II_ADDR	II_DATA				Składnik	
	D3	D2	D1	D0	mnemonik	uwagi
0	bit 3	bit 2	bit 1	bit 0	REJESTR1	Indeks 0
1	bit 7	bit 6	bit 5	bit 4		
2	bit 3	bit 2	bit 1	bit 0	REJESTR2	Indeks 0
3			bit 5	bit 4		
4	bit 3	bit 2	bit 1	bit 0		Indeks 1
5			bit 5	bit 4		
6			bit 1	bit 0	BITY1	Indeks 0
		bit 0			BITY2	Indeks 0
7			bit 1	bit 0	BITY3	Indeks 0
	bit 1	bit 0				Indeks 1
8-11	bit 3	bit 2	bit 1	bit 0	PAMIĘĆ	blok 0
12-15			bit 5	bit 4	PAMIĘĆ	blok 1

u dla
onano
tera i

Alternatywny obraz implementacji tej samej struktury dla 8-bitowej magistrali Interfejsu zamieszczono w tab. 7. W poniższym przykładzie, dzięki szerszej magistrali danych, nastąpiło bardziej ściśle ułożenie składników w obszarze adresowania oraz pełne połączenie bitów w obrębie jednego wektora.

Tabela 7

Implementacja przykładu dla 8-bitowej magistrali Interfejsu

Implementation of an example for 8-bit Interface bus

II_ADDR	II_DATA								Składnik	
	D7	D6	D5	D4	D3	D2	D1	D0	mnemonik	uwagi
0	bit 7	bit 6	bit 5	bit 4	bit 3	bit 2	bit 1	Bit 0	REJESTR1	Indeks 0
1			bit 5	bit 4	bit 3	bit 2	bit 1	Bit 0	REJESTR2	Indeks 0
2			bit 5	bit 4	bit 3	bit 2	bit 1	Bit 0		Indeks 1
3							bit 1	Bit 0	BITY1	Indeks 0
						bit 0			BITY2	Indeks 0
			bit 1	bit 0					BITY3	Indeks 0
	bit 1	bit 0								Indeks 1
4-7			bit 5	bit 4	bit 3	bit 2	bit 1	Bit 0	PAMIĘĆ	blok 0

W tab. 8 zamieszczono rezultaty syntezy dla wybranych przedstawicieli serii FPGA firm Altera i Xilinx. Porównano najważniejsze parametry wykorzystania zasobów oraz uzyskane parametry czasowe na podstawie minimalnego czasu konwersji adresu w układzie FPGA. W tabeli podano parametry implementacji odpowiednio dla 4-bitowej i 8-bitowej magistrali Interfejsu rozdzielone znakiem „/”.

Tabela 8

Synteza przykładu dla 4-bitowej/8-bitowej magistrali Interfejsu

Synthesis of example for 4-bit/8-bit Interface bus

Producent	seria	typ	Użyte bloki logiczne	Użyte bloki pamięci	Czas konwersji adresu [ns]
Altera	Cyclone	EP1C6F256C6	58/45	2/1	17/15
	CycloneII	EP2C5F256C6	58/45	2/1	20/18
	Stratix	EP1S10B672C6	58/45	2/1	19/17
Xilinx	Spartan2	XC2S15CS144-4	66/37	2/1	18/14
	Spartan3	XC3S50TQ144-4	57/39	2/1	15/13
	VirtexII	XC2V40CS144-6	58/37	2/1	13/11

Uzyskano systemowe rozwiązanie behawioralne uniwersalnego interfejsu komunikacyjnego dla układów FPGA, które jest syntezowane w różnych środowiskach projektowych. Wskazują na to zamieszczone w tab. 8 rezultaty syntezy dla kilku różnych rodzin układów FPGA. Są dobrze zbieżne zarówno w zakresie zajętości zasobów jak

TOM
i szy
dyfik

F
zosta
nych.
CMS
układ
mysł
jak ro
fejsu
ka), v
— U
wyko
sjach
RF-G
wane
FERM
N
zależn
mentu
progra
MATI
system
dzie I
SI
wieraj
moduł
korzys
przepr
nym st
Interfe
cyjnej.

- 1. http
- 2. http
- 3. http

i szybkości działania. Zmiana serii FPGA oraz producenta nie wymaga żadnych modyfikacji projektu. Podobnie jak zmiana parametrów magistrali komunikacyjnej.

4. ZAKOŃCZENIE

Efektywność projektowania oraz skuteczność i niezawodność działania Interfejsu została zweryfikowana w praktyce w kilku dużych projektach systemów elektronicznych. Interfejs zastosowano w projekcie Trygera Mionowego RPC w eksperymencie CMS przy akceleratorze LHC (Genewa). Docelowo będzie obsługiwał ponad 5000 układów FPGA umieszczonych na ok. 3000 płyt elektronicznych. Współpracuje z przemysłowym interfejsem VME-BUS, odpornym radiacyjnie układem sterowania CCU, jak również osadzonymi procesorami klasy PC z Ethernetem 100T. Aplikacje Interfejsu zaimplementowano w podprojektach grupy fińskiej (Lapperanta — Politechnika), włoskiej (Bari — Narodowy Instytut Fizyki Jądrowej) oraz polskiej (Warszawa — Uniwersytet Warszawski i Politechnika Warszawska) [16]. Interfejs jest również wykorzystywany w eksperymentach TT2 i VUV-FEL (DESY) w kilku kolejnych wersjach systemu symulacji i sterowania wnęką rezonansową SIMCON oraz do sterowania RF-GUN [24-27]. Systemy sterowania wnęką wyposażone w Interfejs zostały przetestowane i wdrożone do akceleratora liniowego w amerykańskim ośrodku badań jądrowych FERMI-LAB.

Na potrzeby obsługi aparatury sterowanej poprzez Interfejs powstało kilka niezależnych rozwiązań warstwy programistycznej, dostosowanej do wymogów eksperymentu CMS oraz akceleratora TTF2 i VUV-FEL. Wykorzystano do tego celu języki programowania C, C++, Java, jak również specjalizowane środowiska programistyczne MATLAB, DOOCS, XDAQ. Implementacji warstwy programistycznej dokonano w systemach operacyjnych Windows98/2000/XP, Linux i Unix na maszynach w standardzie IBM-PC, SUN i procesorach wbudowanych (ETRAX i POWER-PC/XILINX).

Skuteczność i niezawodność Interfejsu została zweryfikowana przy projektach zawierających od kilkuset do kilku tysięcy składników różnych typów. Kolejne wersje modułów elektronicznych, wyposażone w nowsze generacje układów FPGA mogą wykorzystać wcześniej opracowane projekty. W takim wypadku, projekt wymaga jedynie przeprowadzenia syntezy, natomiast warstwa funkcjonalna i struktura projektu w żadnym stopniu nie muszą być modyfikowane. W konsekwencji, przedstawione rozwiązanie Interfejsu pozwala również znacząco skrócić czas projektowania warstwy komunikacyjnej.

5. BIBLIOGRAFIA

1. <http://www.xilinx.com/> [Xilinx Homepage].
2. <http://www.altera.com/> [Altera Homepage].
3. <http://www.latticesemi.com/> [Lattice Homepage].

4. <http://www.actel.com/> [Actel Homepage].
5. <http://www.quicklogic.com/> [QuickLogic].
6. U. Meyer-Baese: *Digital Signal Processing with Field Programmable Gate Arrays*, Wydanie drugie, Springer, 2004, ISBN: 3540211195.
7. U. Meyer-Baese: *DSP with FPGAs: VHDL Solution Manual*, Wydanie pierwsze, 2004, ISBN: 0975549499.
8. <http://www.altera.com/literature/technology/dsp/dsp-literature.pdf>, „DSP Literature”, Altera Corporation, 2005.
9. K.T. Pożniak, T. Czarski, R. Romaniuk: *Functional Analysis of DSP Blocks in FPGA Chips for Application in TESLA LLRF System*, TESLA Technical Note, 2003-29.
10. A. Athavale, C. Christensen: *High-Speed Serial I/O Made Simple, A Designer's Guide with FPGA Applications*, Preliminary Edition, 2005, Xilinx Connectivity Solutions, PN0402399.
11. K.T. Pożniak, R.S. Romaniuk, W. Jałmuzna, K. Ołowski, K. Perkuszewski, J. Zieliński, K. Kierzkowski: *FPGA Based, Full-Duplex, Multi-Channel, Multi-Gigabit, Optical, Synchronous Data Transceiver for TESLA Technology LLRF Control System*, TESLA Technical Note, 2004-07.
12. R.S. Romaniuk, K.T. Pożniak, G. Wrochna, S. Simrock: *Optoelectronics in TESLA, LHC, and pi-of-the-sky experiments*, Proc. SPIE Vol. 5576, 2005, pp. 299-309.
13. K.T. Pożniak: *Electronics and photonics for high-energy physics experiments*, Proc. SPIE Vol. 5125, 2003, pp. 91-100.
14. K.T. Pożniak: *FPGA based implementation of hardware diagnostic layer for local trigger of BAC calorimeter for ZEUS detector*, Proc. SPIE Vol. 5484, 2004, pp. 193-201.
15. W. Giergusiewicz, W. Koprek, W. Jałmuzna, K.T. Pożniak, R.S. Romaniuk: *FPGA Based, DSP Integrated, 8-Channel SIM-CON, ver. 3.0. Initial Results for 8-Channel Algorithm*, TESLA Technical Note, 2005-14.
16. M.I. Kudła: *RPC Trigger Overview*, RPC Trigger ESR, Warsaw, July 8th, 2003, http://hep.fuw.edu.pl/cms/esr/talks/MK_trigger_overview.pdf.
17. National Instruments Corporation, *LabVIEW — FPGA Module User Manual*, Technical Document, Part Number 370690B-01, 2004.
18. Nallatech, *FUSE System Software User Guide*, NT107-0068V2, Issue 3, 2002.
19. Nallatech, *FUSE Toolbox for MATLAB Product Brief*, Technical Document, <http://www.nallatech.com/mediaLibrary/images/english/2398.pdf>.
20. http://www.altera.com/products/ip/processors/nios/features/nioavalon_bus.html.
21. <http://www.opencores.org/projects.cgi/web/wishbone/wishbone>.
22. K.T. Pożniak, M. Bartoszek, M. Pietrusiński: *Internal Interface for RPC Muon Trigger electronics at CMS experiment*, Proc. SPIE Vol. 5484, 2004, pp. 269-282.
23. K.T. Pożniak: *INTERNAL INTERFACE...* Tesla Note 2005-22, 2005.
24. W. Koprek, P. Kaleta, J. Szewiński, K.T. Pożniak, T. Czarski, R.S. Romaniuk: *Software layer for FPGA-based TESLA cavity control system*, TESLA Report 2004-10, 2004. DESY.
25. K.T. Pożniak, T. Czarski, W. Koprek, R.S. Romaniuk: *SIMCON 2.1. Manual*, Tesla Note 2005-02, 2005, DESY.
26. K.T. Pożniak, T. Czarski, W. Koprek, R.S. Romaniuk: *SIMCON 3.0. Manual*, Tesla Note 2005-05, 2005, DESY.
27. W. Giergusiewicz, W. Koprek, W. Jałmuzna, K.T. Pożniak, R.S. Romaniuk: *FPGA based, DSP board for LLRF 8-Channel SIMCON 3.0 Part I: Hardware*, Proc. SPIE Vol. 5948, 2005, pp. 110-120.

K. POŹNIAK

PARAMETRIZED COMMUNICATION INTERFACE FOR FPGA DEVICES

Summary

Contemporary technology of FPGA provides: hundreds of thousands of logical blocks, up to 10MB of SRAM memory, above 200 fast blocks of DSP, up to 20 nondependent modules for multi-gigabit electrical and optical transmission, all in a single circuit. FPGA circuits are more and more applied in market electronics as well as in large, multichannel, specialized electronic systems. They are equipped in complex communication interfaces. The interfaces are responsible for remote control, system configuration, detailed diagnostics and online monitoring. The paper describes hardware layer of a parameterized, universal communication interface for applications in FPGA circuits. There are debated the methods to automatically create the address and data areas, according to the user's declarations. The methods to standardize the I/O communications are described. Implementation examples are given.

Keywords: Altera, Xilinx, VHDL, parameterization, communication interfaces, behavioral programming

Ś
nego,
optycz
długo:
lające
optycz
zasadn
pagacj
W
dem p
• refra
wspó
mnie
• foton

Wytwarzanie i charakteryzacja światłowodów kapilarnych

RYSZARD S. ROMANIUK¹, JAN DOROSZ²

¹ Instytut Systemów Elektronicznych, Politechnika Warszawska, ul. Nowowiejska 15/19, 00-665 Warszawa

² Katedra Promieniowania Optycznego, Politechnika Białostocka, ul. Wiejska 45c, 15-351 Białystok

Otrzymano 2005.12.08

Autoryzowano 2006.02.03

W pracy opisano podstawowe właściwości transmisyjne światłowodów kapilarnych. Pokazano metody obliczania charakterystyk propagacyjnych. Wytwarzano laboratoryjne próbki światłowodów kapilarnych metodą tyglową i preformową ze szkielek miękkich borokrzemionkowych i wapniowosodowych. Mierzono charakterystyki geometryczne, optyczne i mechaniczne światłowodów. Porównywano obliczenia z pomiarami.

Słowa kluczowe: światłowodory kształtowane, kapilary optyczne, technologia światłowodowa, pomiary światłowodów, technika światłowodowa

1. WSTĘP

Światłowodem kapilarnym nazywamy rodzaj włókna optycznego, na ogół szklanego, wysokokrzemionkowego, posiadającego właściwości niskostratnej transmisji fali optycznej wielomodowej lub jednomodowej na znaczne odległości (w porównaniu z długością fali) i jednocześnie posiadającego właściwości klasycznej kapilary, pozwalające na transport niewielkich ilości materii, współbieżnie lub przeciwbieżnie do fali optycznej. Fala optyczna w światłowodzie kapilarnym może być prowadzona na dwa zasadnicze sposoby: w otworze (próżnia, gaz, ciecz), jeśli spełnione są warunki propagacji niskostratnej, lub w szklanym rdzeniu pierścieniowym wokół otworu.

W otworze kapilary fala optyczna może propagować się kilkoma metodami (modem podstawowym jest HE_{01}):

- refrakcyjną — dla fal rentgenowskich (przy małych kątach propagacji), dla których współczynnik załamania szkła posiada wartość zespoloną i część rzeczywista jest mniejsza od jedności, tzw. całkowite zewnętrzne odbicie;
- fotoniczną — gdy warstwa szkła przy granicy szkło-próżnia jest formowana w

postaci fotonicznej przerwy zabronionej;

- całkowitego odbicia od wewnętrznego pokrycia otworu metalem szlachetnym i warstwą dielektryka.

Propagacja fali optycznej o dużej gęstości mocy w otworze kapilary o niewielkim wymiarze na znaczną odległość stanowi podstawę do aplikacji wykorzystujących oddziaływanie takiej fali z materią. W rdzeniu pierścieniowym fala optyczna jest propagowana metodą refrakcyjną (modem podstawowym jest HE_{11}) — wielomodową i jednomodową. Jeśli rdzeń pierścieniowy jest stosunkowo wąski to pewna część mocy optycznej (fala zanikająca) wnika w otwór kapilary, stanowiąc podstawę do szeregu aplikacji wykorzystujących oddziaływanie silnego gradientu fali zanikającej z materią.

O właściwościach transmisyjnych światłowodu kapilarnego w dużej mierze decyduje proporcja szkło-powietrze (a raczej próżnia lub inne wypełnienie otworu) jaką widzi propagowana fala optyczna. Stąd, kapilary mogą być stosowane do transmisji dużych mocy optycznych, spektroskopii niewielkich ilości współpropagowanej materii, wspomaganie optycznie chromatografii, elektroforezy i reologii, szczególnie w przyszłościowych układach typu MOEMS. Kapilary optyczne posiadają anomalne właściwości dyspersyjne, mogą wykazywać silne optyczne cechy nieliniowe, stanowią elementarny składnik znacznej rodziny włóknowych światłowodów fotonicznych — multi-kapilarnych.

2. POLE ELEKTROMAGNETYCZNE W ŚWIATŁOWODZIE KAPILARNYM

Światłowód kapilarny posiada następujący profil refrakcyjny:

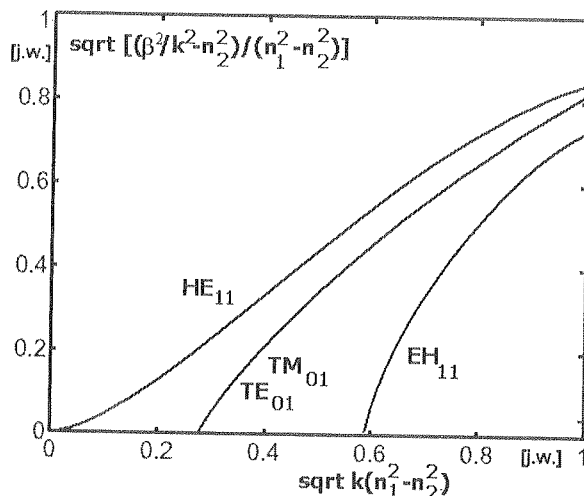
$$n(x, y, z) = n(r) = 1 \text{ dla } r < r_1, n(r) = n_1 \text{ dla } r_1 < r < r_2, n(r) = n_2, \text{ dla } r > r_2 \quad (1)$$

gdzie $r = \sqrt{x^2 + y^2}$ jest odległością od osi symetrii, r_1 jest promieniem otworu kapilarnego, r_2 jest promieniem płaszcza optycznego (do granicy płaszczy — pokrycie), n_1 i n_2 są refrakcjami rdzenia i płaszcza kapilary optycznej. Mody fal E (oraz H) w postaci $E(x, y, z, t) = E_o(x, y) \exp i(\beta z - \omega t)$ spełniają równanie Helmholtza $(\Delta_r^2 + n^2 k^2 - \beta) E_z = 0$, w każdym regionie profilu refrakcyjnego światłowodu n , gdzie β jest stałą propagacji wzdłuż światłowodu, $\omega = CK$ jest częstotliwością kątową pola EM, k jest liczbą falową. Rozwiązania są następujące (dla pola E):

$$\begin{aligned} E_z(r, \phi) &= AI_m(vr) \sin(m\phi + \varphi), r < r_1, \\ &= [BJ_m(ur) + CN_m(ur)] \sin(m\phi + \varphi), r_1 < r < r_2, \\ &= DK_m(wr) \sin(m\phi + \varphi), r > r_2, \end{aligned} \quad (2)$$

gdzie J_m i N_m są funkcjami Bessela pierwszego i drugiego rodzaju rzędu m , I_m i K_m są modyfikowanymi funkcjami Bessela pierwszego i drugiego rodzaju rzędu m , φ jest dowolnym stałym składnikiem fazy. Stałe A, B, C i D są określane z warunków ciągłości składowych pól E_ϕ , E_z na granicach obszarów różnej refrakcji. Argumenty

funkcji Bessela są: $u = \text{sqrt}(k^2 n_1^2 - \beta^2)$ — poprzeczna stała propagacji, $v = \text{sqrt}(\beta^2 - k^2)$, $w = \text{sqrt}(\beta^2 - k^2 n_2^2)$ — poprzeczne stałe tłumienia w obszarze otworu kapilarnego i płaszcza optycznego. Składowe H_z posiadają identyczną postać z czynnikiem ortogonalnym $\cos(m\Phi + \varphi)$. Pozostałe składowe pól określa się ze składowych podłużnych. Równanie własne dyspersyjne otrzymuje się stosując warunki brzegowe do składowych pól [2]. Rozwiązanie równania własnego dla przykładowego światłowodu kapilarnego przedstawiono na rys.1.

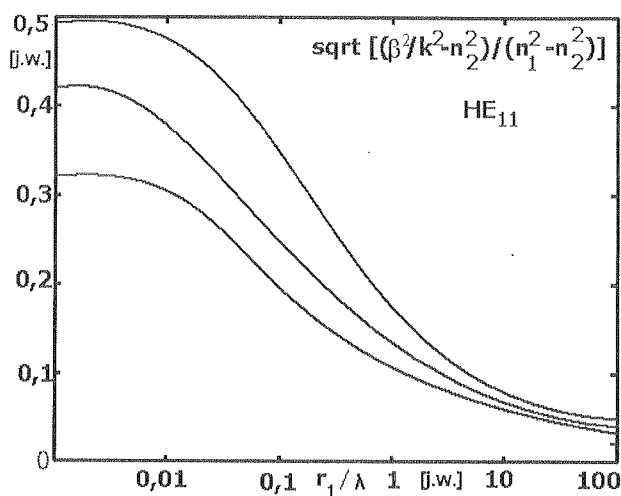


Rys. 1. Charakterystyki dyspersji modowej i odcięcia modowego, w postaci znormalizowanej stałej propagacji jako funkcji znormalizowanej liczby falowej, dla modów najniższego rzędu w światłowodzie kapilarnym: $\lambda = 800\text{nm}$, $2r_1 = 10\mu\text{m}$, $\Delta n = 0,2\%$, $n_2 = 1,50$

Fig. 1. Dispersion characteristics nad modal cut-off, normalized propagation constant as a function of normalized wave number for the lowest order modes in capillary optical fibre: $\lambda = 800\text{nm}$, $2r_1 = 10\mu\text{m}$, $\Delta n = 0,2\%$, $n_2 = 1,50$

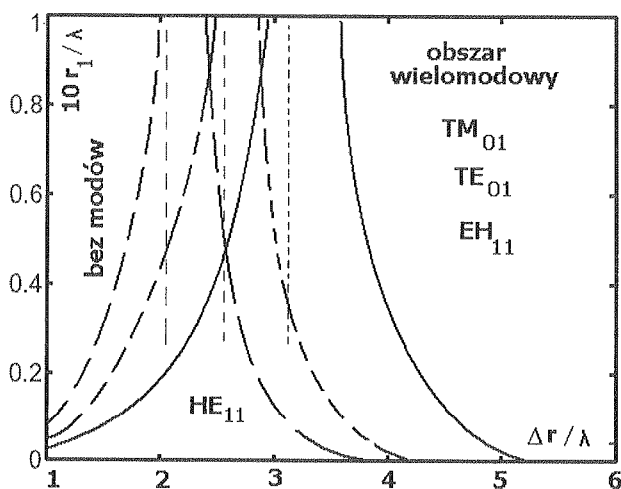
Z podstawowych charakterystyk dyspersyjnych światłowodu kapilarnego, przedstawionych na rys.1. dla przykładowego wyciąganego światłowodu jednomodowego wynikają jego pozostałe właściwości modowe i propagacyjne. Na kolejnych rysunkach przedstawiono znormalizowane obliczone charakterystyki zależności struktury modowej od względnego wymiaru otworu kapilary (w odniesieniu do długości propagowanej fali). Obliczenia przeprowadzono w środowisku MatLab (toolbox EM). Rys. 2. przedstawia charakterystykę modu HE_{11} (dla azymutalnej liczby modowej $m=1$) dla trzech różnych światłowodów. Najwyższa krzywa jest dla światłowodu z rys. 1.

Zmiennymi niezależnymi w obliczeniach charakterystyk światłowodu kapilarnego są r_1/λ , r_2/λ , β/k , n_1 , n_2 . Gdy średnica otworu światłowodu kapilarnego jest duża, struktura modowa jest podobna do jednowymiarowego, asymetrycznego światłowodu planarnego (mody TE i TM). Zmniejszanie wartości r_1 , dla stałej wartości grubości rdzenia pierścieniowego, prowadzi do propagacji modu podstawowego HE_{11} . Obszar



Rys. 2. Zmiana struktury modowej modu podstawowego HE_{11} światłowodu kapilarnego w funkcji znormalizowanej średnicy otworu kapilarnego

Fig. 2. Change of HE_{11} modal structure in capillary optical fibre as function of normalized capillary hole



Rys. 3. Obszary różnej struktury modowej w światłowodzie kapilarnym w funkcji względnych wymiarów otworu kapilarnego oraz grubości pierścienia rdzeniowego $\Delta r = r_2 - r_1$ dla trzech różnych światłowodów. Charakterystyka środkowa jest dla światłowodu z rys. 1

Fig. 3. Areas of different modal structure in capillary optical fiber as function of relative dimensions of capillary hole and the thickness of ring optical core $\Delta r = r_2 - r_1$ for three different optical fibres.

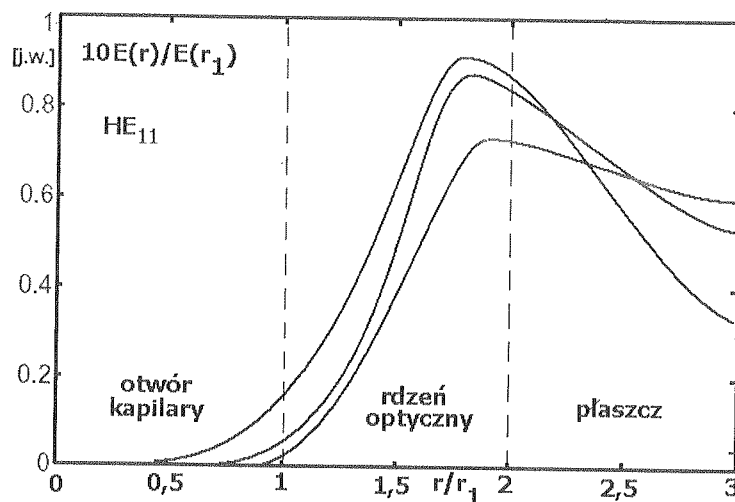
Middle characteristic is for fibre from fig. 1

Rys. 4. W otworze

Fig. 4. Re

jednomo
tłowodów
warunek
kresu gr
nie zależ
stałych p
światłow
stawiono
od wymi

Podst
wykonane
mysłu tek
wej. Zakł
90-tych, w
wań poza
niono taki
fazie chłoc
cji, złożon



Rys. 4. Względne natężenie pola elektrycznego dla modu HE_{11} w funkcji znormalizowanego promienia otworu kapilary, dla trzech różnych światłowodów. Krzywa środkowa jest dla światłowodu z rys.1

Fig. 4. Relative intensity of E field for HE_{11} mode as function of normalized radius of capillary hole, for three fibres. The middle curve is for fibre from fig.1

jednomodowości światłowodu kapilarnego przedstawiono na rys.3, dla trzech światłowodów o różnych parametrach. Asymptotyczna prosta linia przerywana oznacza warunek odcięcia modu TE_1 dla odpowiedniego światłowodu planarnego. Dla tego zakresu grubości rdzenia w pobliżu odcięcia modowego, warunki propagacji praktycznie nie zależą od wielkości otworu kapilary. Rozwiązanie problemu własnego i obliczenie stałych propagacji pozwala na znalezienie rozkładu pola E w przekroju poprzecznym światłowodu kapilarnego. Rozkład pola dla trzech analizowanych światłowodów przedstawiono na rys. 4. Głębokość wnikania pola w otwór kapilarny w istotny sposób zależy od wymienionych parametrów włókna optycznego.

3. WYCIĄGANIE ŚWIATŁOWODÓW KAPILARNYCH

Podstawowe studia teoretyczne nad stacjonarnym wyciąganiem włókien zostały wykonane na przełomie lat 60 i 70 [4], głównie przez zespół J.Pearsona dla przemysłu tekstylnego. Wyniki te z powodzeniem adaptowano dla techniki światłowodowej. Zakładano model wyciągania słabo-perturbacyjny (idealizowany). Na początku lat 90-tych, wraz z silnym rozwojem specjalizacji światłowodów, szczególnie do zastosowań poza-telekomunikacyjnych, w modelu wyciągania włókna optycznego uwzględniono takie czynniki jak: transfer ciepła z wiązki laserowej, stożkowanie, zaburzenia w fazie chłodzenia, kształtowanie początkowej kropki, uwzględnienie sił inercji i grawitacji, złożoną strukturę preformy, poziom lepkości szkła poprzez wartość współczynnika

Rayleigha, transport ciepła, silne niestabilności mechaniczne [5]. Analizie poddano także procesy wyciągania szklanego włókna i światłowodu kapilarnego [6], łącznie z analizą deformacji otworu kapilarnego.

3.1. ROZWIĄZANIE RÓWNAŃ NAVIERA-STOKESA DLA ŚWIATŁOWODU KAPILARNEGO

Przekształcone równania Naviera-Stokesa i dyfuzyjno konwekcyjne dla warunków kapilary szklanej przedstawiono poniżej, łącznie z rysunkiem przedstawiającym podstawową geometrię wyciągania światłowodu kapilarnego. Wyjściowym założeniem jest, że średnica kapilary jest znacznie mniejsza od charakterystycznej długości menisku wypływowego $r \ll L$. Równania te wynikają z praw ciągłości przepływu, równowagi momentów pędu oraz bilansu energii.

$$\rho(r_2^2 - r_1^2)[v_t + v v_z - g] = [3\mu(r_2^2 - r_1^2)v_z + \xi(r_1 + r_2)]_z$$

$$(r_1^2)_t + (r_1^2 v)_z = \frac{p_0 r_1^2 r_2^2 - \xi r_1 r_2 (r_1 + r_2)}{\mu(r_2^2 - r_1^2)}$$

$$(r_2^2)_t + (r_2^2 v)_z = \frac{p_0 r_1^2 r_2^2 - \xi r_1 r_2 (r_1 + r_2)}{\mu(r_2^2 - r_1^2)}$$

$$\frac{r_2^2 - r_1^2}{2} [\rho c_p (T_t + v T_z) - k (T_z)_z - \sigma \varepsilon' (T_a^4 - T^4)] = r_2 h (T_a - T)$$

Znormalizowane rozwiązania równań Naviera-Stokesa w postaci rodzin krzywych parametrycznych przedstawiono na rys. 6. Parametrami są wielkości technologiczne: prędkość podawania preformy, prędkość wyciągania włókna optycznego, średnica włókna, proporcje wymiarowe w kapilarze, temperatura procesu.

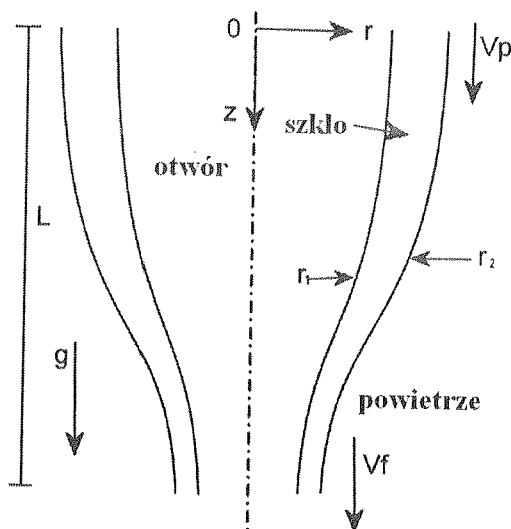
Analiza powyższych równań prowadzi do szeregu wniosków dotyczących warunków technologicznych kształtowania kapilary światłowodowej.

Względna waga składników inercyjnego, grawitacyjnego i napięcia powierzchniowego jest wyrażona odpowiednio poprzez następujące parametry bezwymiarowe: $LV_f \rho / \mu$, $gL^2 \rho / \mu V_f$, $\xi L / \mu R_f V_f$, gdzie L — charakterystyczna długość menisku o wysokiej temperaturze (10–50 mm), V_f — charakterystyczna prędkość wyciągania włókna optycznego (10–300 m/min), R_f — wymiar otworu kapilarnego (0,1–10 μm dla światłowodu jednomodowego, 30–500 μm dla światłowodu wielomodowego, 0,1–2 mm dla kapilar dyskretnych). Typowe wartości parametrów dla wysokokrzemionkowych szkiele wieloskładnikowych są: $\rho = 2100\text{--}2600 \text{ kg/m}^3$ (2200 kg/m^3), $\mu = 10^4\text{--}10^5$ Pas,

Rys.
jedna
funkcj
k
wyp
róż
pojem
stała S

Fig. 5.
param
deriva
radius,
— pre
exte

$\xi = 0$
dać, że
inercy
napięc
tworze
szczeg
uprosz
Ut
współc

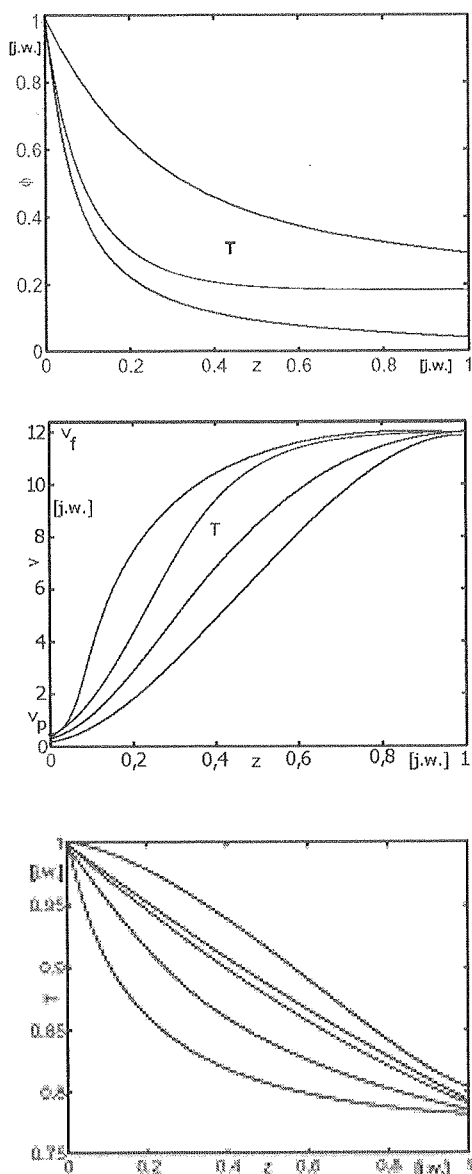


Rys. 5. Układ geometryczny i parametry menisku wyciągowego włókna kapilarnego. Oznaczenia są jednakowe dla rys.5 i wzorów Naviera-Stokesa dla kapilary, indeksy dolne oznaczają pochodną danej funkcji względem zmiennej niezależnej reprezentowanej przez indeks, t — czas, z — dystans wzdłuż osi kapilary, r_1 i r_2 są promieniem wewnętrznym i zewnętrznym kapilary, L — długość menisku wypływowego, v — prędkość, ρ — gęstość, g — przyspieszenie grawitacyjne, μ — lepkość, p_o — różnica ciśnień pomiędzy wnętrzem i zewnętrzem kapilary, ξ — napięcie powierzchniowe, c_p — pojemność cieplna, T — temperatura, T_a — temperatura zewnętrzna, k — przewodność cieplna, σ — stała Stefana-Boltzmana, ε' — stała materiałowa związana z emisyjnością, h — współczynnik transferu ciepła

Fig. 5. Geometrical set-up and parameters of outflow meniscus for pulling of capillary optical fibre. The parameters are identical for fig.5 and the following Navier-Stokes flow equations. Lower indexes are derivatives against this argument. t — time, z — axial distance, r_1 and r_2 internal and external capillary radius, L — length of outflow meniscus, v — velocity, ρ — density, g — gravitation, μ — viscosity, p_o — pressure difference in capillary, ξ — surface tension, c_p — heat capacity, T — temperature, T_a — external temperature, k — heat conductivity, σ — Stefan-Boltzman constant, ε' — material constant combined with emissivity, h — heat transfer coefficient

$\xi = 0,3\text{N/m}$. Z porównania obliczonych wartości parametrów bezwymiarowych wiadać, że składniki lepkościowy i grawitacyjny są zawsze ważne, podczas gdy składnik inercyjny może być pominięty. Ponadto, w obszarze menisku wypływowego, efekty napięcia powierzchniowego i lepkości działają w przeciwnych kierunkach. Warunki tworzenia proporcji geometrycznych światłowodowego włókna kapilarnego zależy od szczegółów proporcji pomiędzy wymienionymi składnikami. Poniżej przedstawiono w uproszczeniu dyskusję wielkości parametrów.

Utrzymanie ciśnienia w otworze kapilarnym jest charakteryzowane przez kolejny współczynnik bezwymiarowy $Lp_o/V_{f\mu}$, gdzie p_o — różnica ciśnień między kapilarą i



Rys. 6. Parametryzowane, rodziny stacjonarnych rozwiązań względnych równań Naviera-Stokesa dla włókna szklanego dla strefy gorącej (zwanej meniskiem wyciągowym). a) funkcyjna ewolucja prędkości podawania preformy w prędkość wyciągania włókna optycznego, b) ewolucja średnicy menisku od preformy do włókna, c) funkcja rozkładu temperatury w strefie menisku

Fig. 6. Parameterized families of relative stationary solutions of N-S equations for optical fibre and for the hot outflow meniscus. a) functional evolution of the shift velocity for the perform into the velocity of fibre pulling, b) evolution of the outflow meniscus radius from perform to fibre, c) temperature distribution in meniscus

otocze
ma odn
wego p
W

pilarne
Układ g

Prz
zachod
stopień
oraz le
szybkos
Poniew
larnego
mała le
długość
na założ
grawita

Jeś
padku i
otworze
rącej je

Jeś
rze, wó
nieduży
go, czu
gdzie r
formy.

technol
układ te
kach ka
preform
danego
niewiele
lepkość
wymiar
preform
tur wyc
preform
wartość
gania ś
 $V_f = 19$

otoczeniem (do ok. 10kPa). Jeśli efekt podwyższonego ciśnienia w otworze kapilarnym ma odnieść widoczny skutek technologiczny to, wartość współczynnika bezwymiarowego powinna być ok. jedności.

W przypadku braku napięcia powierzchniowego oraz różnicowego ciśnienia kapilarnego, otwór kapilarny nie podlega kolapsowi lub ekspansji w czasie wyciągania. Układ geometryczny kapilary jest zachowywany w proporcjach z preformy.

Przy założeniu małego lecz zauważalnego wpływu napięcia powierzchniowego (co zachodzi dla wysokiej temperatury procesu) dominują zjawiska lepkościowe. Wówczas, stopień zamykania (kolapsu) otworów zależy od stosunku napięcia powierzchniowego oraz lepkości ξ/μ . W tych warunkach, proces zamykania otworu jest bardziej czuły na szybkość wprowadzania preformy do pieca niż na szybkość wyciągania światłowodu. Ponieważ lepkość jest silną funkcją temperatury, więc proces kolapsu otworu kapilarnego także silnie zależy od temperatury. Zamykaniu otworów kapilarnych sprzyja mała lepkość (wysoka temperatura procesu), powolne podawanie preformy oraz duża długość obszaru menisku (strefy gorącej). Uproszczenie powyższych rozważań polega na założeniu izotermiczności całego obszaru menisku, oraz pominięciu efektów inercji, grawitacji oraz nadciśnienia we wnętrzu kapilary.

Jeśli siły inercji, grawitacji oraz napięcia powierzchniowego są pomijane w przypadku izotermicznym, wówczas czułość procesu technologicznego na nadciśnienie w otworze kapilarnym jest znacznie zwiększona. Pod warunkiem, że długość strefy gorącej jest duża, lepkość szkła jest mała oraz preforma jest podawana wolno.

Jeśli brane są pod uwagę napięcie powierzchniowe oraz nadciśnienie w kapilarze, wówczas można ocenić czułość procesu na to nadciśnienie. Zakładając względnie nieduży otwór kapilarny, pomijając inercję i grawitację, dla przypadku izotermicznego, czułość nadciśnieniową można wyrazić zależnością [4] $S = L\xi/\mu r_{ID} V_p \log(V_f/V_p)$, gdzie $r_{ID} = r_1(0)$ — wewnętrzny promień preformy, V_p — prędkość podawania preformy. Gdy $S \ll 1$ to nadciśnienie w kapilarze może być skutecznym parametrem technologicznym kształtującym proporcje wymiarowe: szkło-powietrze. Gdy $S \gg 1$ to układ technologiczny jest bardzo czuły na nadciśnienie i wyciąganie w takich warunkach kapilary jest niestabilne. Parametr S jest bardziej czuły na prędkość podawania preformy niż na prędkość wyciągania światłowodu. Ponieważ długość strefy gorącej dla danego procesu technologicznego jest stała, oraz napięcie powierzchniowe zmienia się niewiele dla szkieł wysokokrzemionkowych, to współczynnik czułości zależy tylko od lepkości szkła (tzn. od temperatury procesu wyciągania światłowodu), początkowego wymiaru otworu w preformie, prędkości podawania preformy oraz stosunku prędkości preformy do światłowodu V_f/V_p . Współczynnik czułości wzrasta dla niższych temperatur wyciągania włókna, mniejszych wymiarów otworu, mniejszych prędkości podawania preformy oraz mniejszych wartości stosunku prędkości V_f/V_p . Przykładowe obliczone wartości parametru czułości wymiarów otworu kapilary na parametry procesu wyciągania światłowodu zebrano w tabeli 1, dla następujących danych: $r_1(0) = 1,5$ mm, $V_f = 190$ m/min, $V_p = 3$ mm/min, $T = 950^\circ\text{C}$.

Tabela 1

Obliczona czułość technologiczna wyciągania kapilary w funkcji prędkości podawania preformy V_p , prędkości wyciągania włókna V_f oraz temperatury T

Calculated technological sensitivity of capillary optical fibre pulling as function of perform shift V_p , fibre pulling V_f and temperature T

$V_{f,p}$	S		T [°C]	S
0,2	1,9		T_o-100	0,1
1	0,4		T_o	0,4
1,6	0,2		T_o+100	2,6

Powyższe wnioski poczyniono przy założeniu izotermiczności procesu technologicznego zachodzącego w strefie menisku. W rzeczywistości niejednorodny rozkład radialny i osiowy temperatury dotyczy w największym stopniu lepkości szkła.

Głównym parametrem decydującym o zamykaniu otworu kapilarnego jest stosunek napięcia powierzchniowego do lepkości szkła. Napięcie powierzchniowe szkieł wysokokrzemionkowych jest słabo zależne od temperatury. Lepkość szkieł wysokokrzemiankowych i czystego szkła krzemionkowego silnie zależy od temperatury. Stąd, temperatura procesu wyciągania światłowodu kapilarnego jest podstawowym czynnikiem decydującym o kolapsie otworu kapilarnego. Podobnie, prędkość podawania preformy bardziej decyduje o procesie zamykania otworów niż prędkość wyciągania włókna.

Wymiary względne szklanej kapilary optycznej zależą od szybkości zmian grubości ściany kapilary, w obszarze menisku, wzdłuż kierunku wyciągania włókna. Ta szybkość zmian zależy istotnie od kształtu menisku wpływowego, czyli od krzywizny wewnętrznej i zewnętrznej powierzchni tworzonej kapilary. W przypadku optycznego włókna szklanego pełnego menisk wyciągowy zależy od większości parametrów procesu wyciągania jak: prędkość podawania preformy, temperatura wyciągania, wymiary preformy i światłowodu, prędkość wyciągania włókna optycznego.

3.2. EKSPERYMENTY TECHNOLOGICZNE ZE ŚWIATŁOWODAMI KAPILARNYMI

W praktyce wieża wyciągowa włókna optycznego klasycznego i kapilarnego wygląda identycznie. Fotografię wieży (Katedra Promieniowania Optycznego, Politechniki Białostockiej) oraz schemat blokowy związanych z nią podstawowych urządzeń przedstawiono na rys.7 [6].

Parametry dotyczące wytwarzania kapilar można podzielić na dwie grupy: parametry procesu wyciągania, możliwe do zmiany w czasie procesu, oraz parametry strukturalne, geometryczne związane ze strukturą i wymiarami preformy i światłowodu. Najważniejsze parametry procesu to: temperatura wyciągania, prędkość podawania preformy, nadciśnienie preformy kapilarnej. Parametry kapilary, które można zmieniać w czasie procesu są: średnica zewnętrzna (D_z) oraz średnica wewnętrzna (D_w) określające grubość ścianki kapilarnej. Stosunek obu wymiarów można zmieniać w czasie wyciągania poprzez wpływanie na procesy kolapsu lub ekspansji otworu kapilarne-

go. Gdy stosunek wymiarów D_z/D_w we włóknie jest większy niż w preformie, to w czasie procesu wyciągania doszło do częściowego kolapsu. W przeciwnym przypadku doszło do ekspansji otworu. Pomiary zmian proporcji otworu kapilarnego dla różnych preform przedstawiono na rys.8. W kilku przypadkach, wraz ze wzrostem temperatury (przy utrzymywaniu stałego wymiaru zewnętrznego kapilary) otwór wykazywał tendencję zamykania. Podczas wytwarzania kapilary podstawowym problemem jest kontrola i utrzymywanie stałej grubości ścianki d . Wymiary te są związane zależnością $D_z/D_w = D_z/(D_z + 2d) = (D_w + 2d)/D_w$. Zapobieganie kolapsowi kapilary polega na maksymalizacji sił lepkości w czasie wyciągania poprzez obniżenie temperatury procesu. Możliwość kontroli grubości ścianki kapilary pozwala na wytwarzanie szeregu wymiarowego kapilar z jednej preformy.

Tabela 2

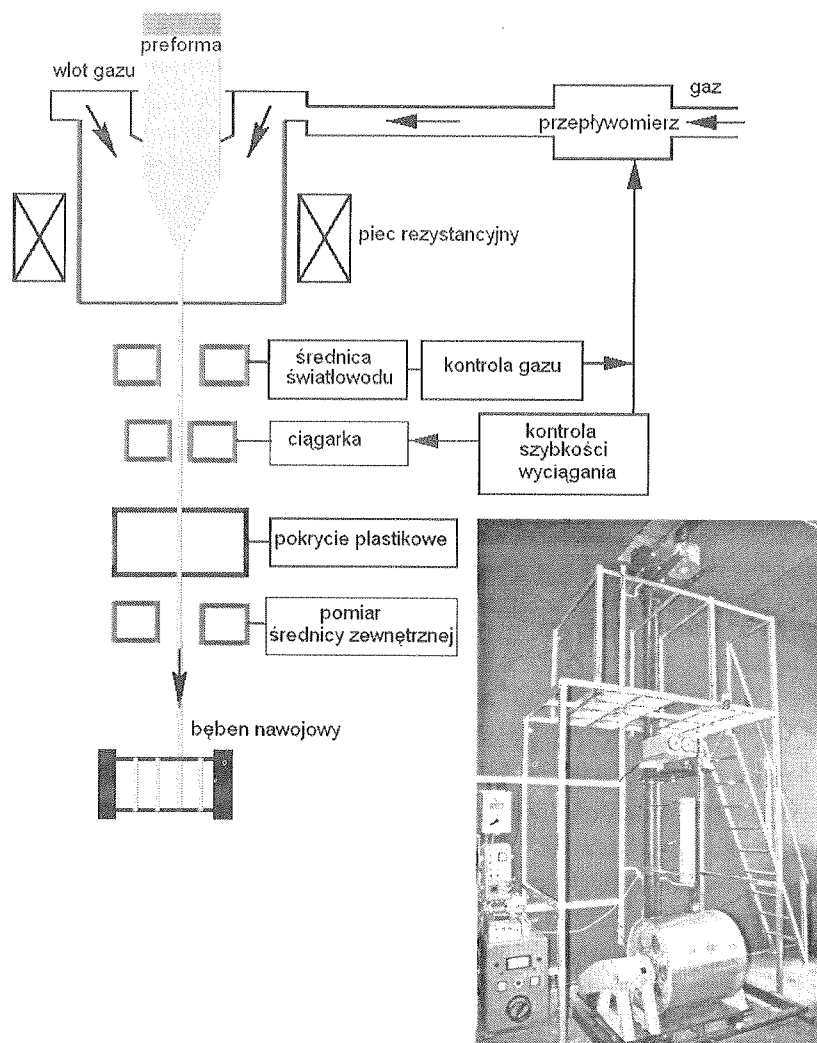
Pomiarowe porównanie zmian proporcji wymiarów preformy i kapilary

Comparison of measurements in the changes of geometrical proportions in perform and capillary

preforma D_z/D_w [mm]	preforma		kapilara	
	wart.	odchylenie	wart.	odchylenie
	średnia D_z/D_w	standardowe [%]	średnia D_z/D_w	standardowe [%]
8x1,1	1,075	0,01	2,3	0,3
34x1,5	1,0075	0,02	1,15	1,1

Wzdłużna stabilność wymiaru poprzecznego światłowodu kapilarnego jest równie ważnym parametrem jak stała wartość stosunku wymiaru zewnętrznego do otworu. Mierzono wzdłużną stabilność wymiarów kapilary i przykładowe wyniki dla światłowodu o wymiarach nominalnych $125\mu\text{m}$ przedstawiono na rys. 9. W celu otrzymania dobrej jakości światłowodu kapilarnego zmiany średnicy zewnętrznej powinny być minimalizowane. Źródła zmian średnicy zewnętrznej są: zmiany średnicy preformy, zmienne parametry procesu wyciągania, drgania mechaniczne układu wieży wyciągowej, stabilność temperatury pieca, parametry laminarnego przepływu gazu izolacyjnego w piecu, nadciśnienie w kapilarze. Zgodnie z zasadą zachowania przepływu, zmiana wymiaru preformy o średnicy 34mm i długości 1mm będzie rozciągnięta na długość kilkuset mm dla światłowodu kapilarnego o średnicy $125\mu\text{m}$. Zmiana średnicy preformy wpływa na kształt menisku wypływowego szkła i rozplływ gazu wokół menisku a przez to zmienia dynamikę procesu wyciągania.

Prędkość podawania preformy światłowodu kapilarnego zmienia temperaturę w piecu w okolicy menisku. Im szybciej podawana jest preforma tym krócej przebywa w strefie gorącej pieca. Zmianie ulega rozkład lepkości w menisku co wpływa na mechanizmy kolapsu. Szybkość podawania preformy, podobnie jak temperatura, jest jednym z parametrów kontroli proporcji wymiarowych światłowodu kapilarnego. Nie



Rys. 7. Fotografia wieży wyciągowej do wytwarzania światłowodów oraz schemat urządzeń technologicznych

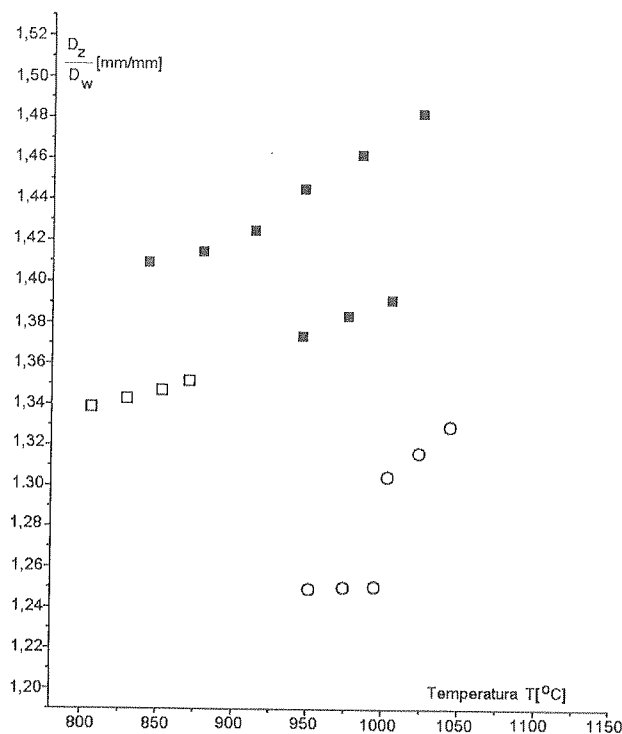
Fig. 7. Photo of the pulling tower for optical fibres and block diagram of technological equipment

Rys. 8. Zmiana funkcji

Fig. 8. Czynność funkcji

jest jedna...
Wymiary...
z czym p...
względne...
Jeśli poda...
lapsowi. W...
zmiany w...

Na rys...
rów zewn...
parametry...
szkieł świ...
my, więk...
większych...
czułe na p...
typowych

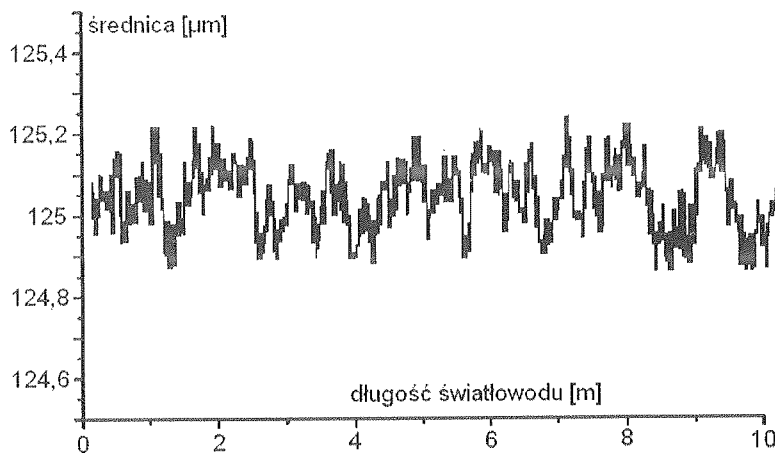


Rys. 8. Zmiana stosunku wymiarów zewnętrznego D_z do wewnętrznego D_w światłowodu kapilarnego w funkcji temperatury wyciągania obrazująca procesy kolapsu lub ekspansji otworu kapilarnego

Fig. 8. Change in the relation between outside D_z and inside D_w dimensions of optical capillary as function of temperature. This characteristic images the collapse or expansion processes of the capillary hole

jest jednak parametrem skutecznym ze względu na duże stałe czasowe procesów zmian. Wymiary otworu kapilary silnie zależą od prędkości podawania preformy, w związku z czym prędkość podawania musi być stabilna. Na rys. 10 przedstawiono zmierzone względne zależności wymiaru otworu kapilarnego od prędkości podawania preformy. Jeśli podawanie preformy jest bardzo wolne, to otwór kapilarny ulega całkowitemu kolapsowi. W takich warunkach, małe zmiany w prędkości podawania powodują znaczne zmiany w wymiarze otworu kapilary.

Na rys. 11 przedstawiono wyniki pomiarów i porównanie z obliczeniami dla wymiarów zewnętrznego i wewnętrznego kapilary. Dane do obliczeń stanowiły następujące parametry: długość strefy grzejnej pieca, parametry wyciągania, dane fizykochemiczne szkieł światłowodowych. Wyższe temperatury, mniejsze prędkości podawania preformy, większe prędkości wyciągania prowadzą do mniejszych rozmiarów kapilary. Dla większych prędkości podawania preformy, ostateczne wymiary kapilary są bardziej czułe na prędkość wyciągania niż dla większych prędkości podawania preformy. Dla typowych stosowanych wartości procesu wyciągania zjawisko zamykania otworu kapi-



Rys. 9. Zmiany średnicy zewnętrznej światłowodu kapilarnego. Preforma $34 \times 1,5$, $T = 1050^\circ\text{C}$, $V_f = 11\text{m/min}$, $D_{zf} = 125\mu\text{m}$. Charakterystyczna częstotliwość zmian wymiarów światłowodu jest znacznie większa niż dla reformy

Fig. 9. Changes of external dimensions of capillary optical fibre. Pereform $34 \times 1,5$, $T = 1050^\circ\text{C}$, $V_f = 11\text{m/min}$, $D_{zf} = 125\mu\text{m}$. Characteristic frequency of dimension change for fibre is much higher than for the preform

larnego występuje relatywnie słabo. Zjawisko to jest silną funkcją prędkości podawania preformy.

Na rys. 12. przedstawiono dwie fotografie wyciąganych światłowodów kapilarnych o różnych proporcjach wymiarowych szkło-powietrze. Kapilary wyciągnięto zgodnie z metodami kształtowania geometrii przedstawionymi powyżej. Przedstawiona metoda kształtowania geometrii jest praktycznym rozwiązaniem problemu produkcji różnych rozwiązań kapilar stosując jedną rodzinę wymiarową preform. Prowadzi to do znacznych oszczędności.

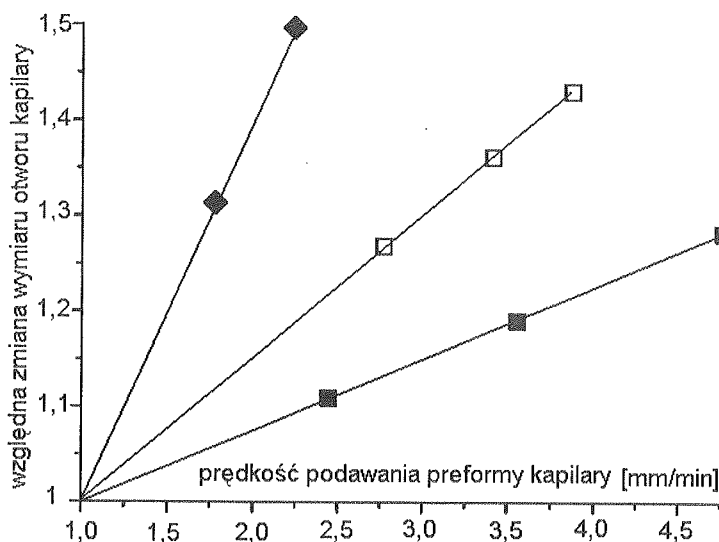
4. WŁAŚCIWOŚCI MECHANICZNE ŚWIATŁOWODÓW KAPILARNYCH

Światłowody wytwarzane z syntetycznej krzemionki mają teoretycznie bardzo dużą wytrzymałość mechaniczną, powyżej 10GPa, nawet większą od porównywalnych włókien stalowych [5]. Mierzone w praktyce wytrzymałości są znacznie mniejsze, rzędu kilku GPa. Jest to spowodowane obecnością w szkłe mikro-szczelin objętościowych i powierzchniowych. Metody technologiczne zwiększania wytrzymałości światłowodu polegają na eliminacji bądź zmniejszeniu takich szczelin. Zerwanie (złamanie) włókna jest inicjowane na szczelinie. Wytrzymałość dynamiczna światłowodu odpowiada sile wymaganej do nagłego zerwania włókna. Wytrzymałość statyczna jest związana z zastosowaniem obciążenia stałego światłowodu przez długi okres czasu. Dynamiczny proces zerwania włókna szklanego (materiał kruchy) ma charakter statystyczny. Ocena

Rys.

Fig. 10

proces
nej rep
w nast
dobień
 S_o — c
mierzo
w skali
jest to
jącego,
optyczn
kien op
od włók
wytrzym
mniej w
Wy
Zjawisk
W szkłe
Szczelin
cząstecz
kumulac
jest opis

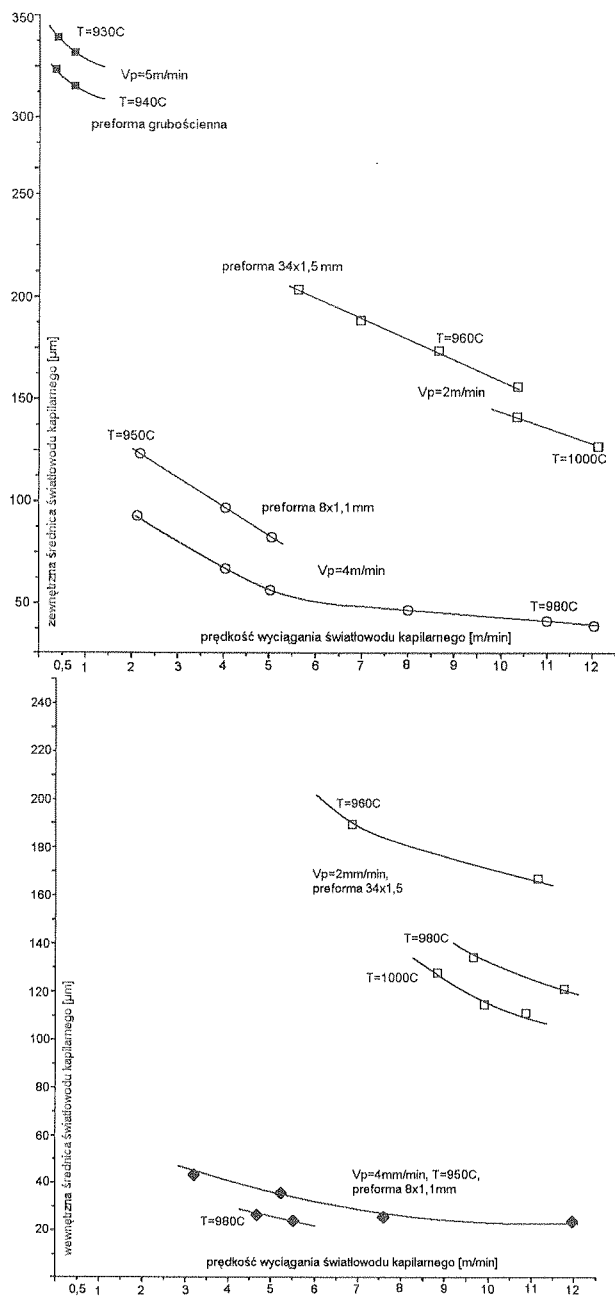


Rys. 10. Względny wzrost wymiaru otworu kapilarnego ze wzrostem prędkości podawania preformy światłowodu kapilarnego dla trzech różnych reform

Fig. 10. Relative increase in dimensions of the capillary hole with the increase in the speed of preform feed velocity, for three different preforms

procesu dokonywana jest przez pomiar dużej ilości próbek w celu otrzymania adekwatnej reprezentacji rozkładu mikro-szczelin w szkłe. Funkcja Weibulla jest przedstawiana w następującej postaci: $\ln \ln(1 - F) - 1 = w \ln(S/S_o) + \ln(L/L_o)$, gdzie: F — prawdopodobieństwo zerwania włókna, w — nachylenie Weibulla, S — naprężenie zrywające, S_o — charakterystyczne naprężenie Weibulla (naprężenie dla $F \approx 0,632$), L — długość mierzonej próbki, L_o — jednostkowa długość próbki mierzonej. Na wykresie Weibulla, w skali podwójnie logarytmicznej, prawie pionowe nachylenie krzywej (w przybliżeniu jest to linia prosta) pomiarów (duża wartość w) dla dużej wartości naprężenia zrywającego, oznacza bardzo skupiony rozkład wytrzymałości i wysokiej jakości włókno optyczne. Na rys. 13 przedstawiono wyniki pomiarów zrywania kilku rodzajów włókien optycznych [7]. Włókna z czystej krzemionki są kilkukrotnie bardziej wytrzymałe od włókien ze szkła miękkich. Włókna kapilarne krzemionkowe są prawie tak samo wytrzymałe jak światłowody klasyczne. Kapilary światłowodowe, szklane, miękkie są mniej wytrzymałe od światłowodów miękkich.

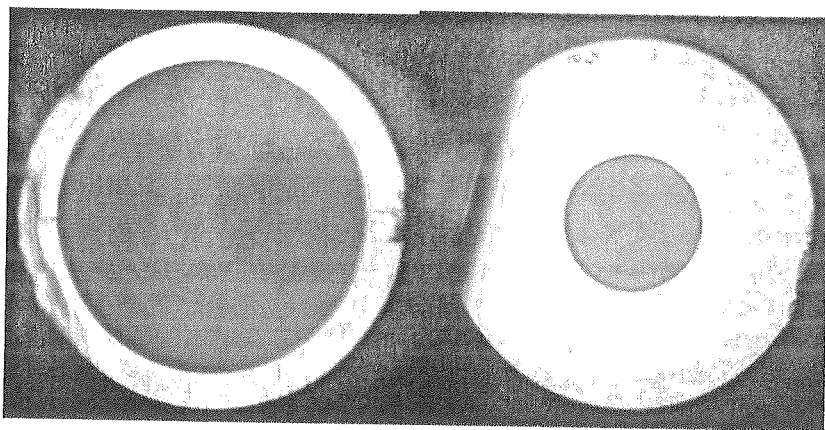
Wytrzymałość mechaniczna światłowodów szklanych ulega degradacji w czasie. Zjawisko określane jest terminem zmęczenia statycznego lub korozją naprężeniową. W szkłe mechanizm ten jest związany z propagacją powierzchniowej mikro-szczeliny. Szczelina propaguje pod wpływem naprężenia, wilgotności (oddziaływania jonowe z cząsteczkami wody), temperatury i czasu. Na szczycie szczeliny naprężenia podlegają kumulacji. Wilgotność i temperatura przyspieszają proces korozji. Mechanizm korozji jest opisany równaniem empirycznym Charlesa [2], $\log(t_z) = Z_s \log(s_{z1}) - \log(s_f)$, gdzie



Rys. 11. Wyniki eksperymentalne (zaznaczone punkty pomiarowe) oraz obliczenia z rozwiązań równań N-S — rys.6 (linie ciągłe) na zewnętrzną i wewnętrzną średnicę kapilary w funkcji prędkości wyciągania światłowodu dla różnych preform, temperatur wyciągania i prędkości podawania preformy V_p

Fig. 11. Experimental data (measured points) and calculations from N-S equations — fig.6. (solid curves) for the external and internal diameter of optical capillary as function of fibre pulling velocity for different preforms, pulling temperatures and preform feeds V_p

t_z — czas
jednostek
zmęczenia
optyczne
wiane ja
wypadku
Wartości
100 dla n
metryczne
Dla włók
szkieł mię
Pomiaria
twierdza o
naprężenia
we włókni
część po
spełnia mi
100kpsi. W
optycznego
wanego cz
kalibrowan
nych włók
[13]. Naprę
skręcające.
powierzchn
ale zabezpi

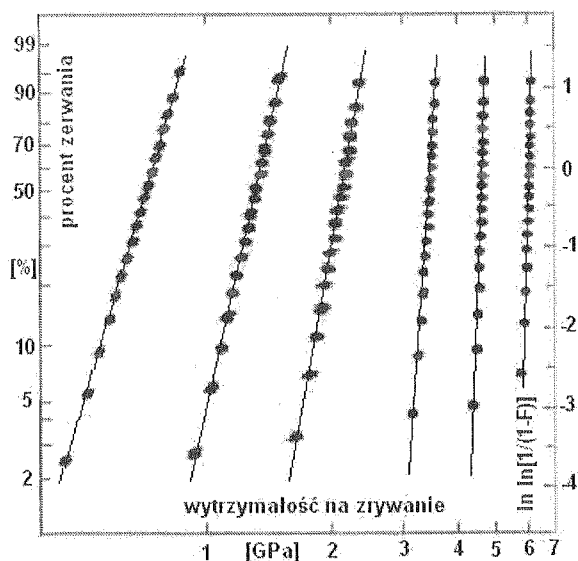


Rys. 12. Światłowody kapilarne wyciągane metodą preformową

Fig. 12. Photos of cross-sections of pulled capillary optical fibres

t_z — czas do zerwania w sekundach, Z_s — parametr zmęczenia statycznego, s_{z1} — jednosekundowe naprężenie zerwania, s_z — naprężenie zerwania w [kpsi]. Parametr zmęczenia statycznego Z_s jest ważną stałą zależną od warunków wytwarzania włókna optycznego oraz materiału. Wyniki zmęczenia statycznego światłowodu są przedstawiane jako logarytm czasu do zerwania w funkcji logarytmu naprężenia. W takim wypadku, krzywa zmęczenia statycznego jest linią prostą o nachyleniu równym Z_s . Wartości parametru Z_s są w zakresie ok. 10 dla szkieł borokrzemionkowych do ponad 100 dla najbardziej wytrzymałych światłowodów krzemionkowych o konstrukcji hermetycznej. Typowe wartości dla światłowodów telekomunikacyjnych wynoszą 20-30. Dla włókien optycznych ze szkieł wieloskładnikowych $Z_s = 5 - 10$ a dla kapilar ze szkieł miękkich określano tą wielkość w zakresie 3-7.

Pomiarowym parametrem wytrzymałości jest tzw. ciągły test wytrzymałości. Potwierdza on wytrzymałość całej długości włókna dla określonego testowego poziomu naprężenia. Ciągły test wytrzymałości jest stosowany do odfiltrowania słabszych miejsc we włóknie optycznym. W sensie mikroskopowym jest to filtracja mikro-szczelin (najczęściej powierzchniowych) o progowej wartości wymiaru. Test zapewnia, że włókno spełnia minimum wytrzymałości. Typową wartością progową testu jest 0,632GPa lub 100kpsi. Wartość testu wytrzymałościowego dobierana jest do rodzaju aplikacji włókna optycznego. Wynik testu daje gwarancje wytrzymałości włókna i/lub jego przewidywanego czasu życia. Test wytrzymałości może być wykonany poprzez zastosowanie kalibrowanego wygięcia włókna lub kalibrowanej siły rozciągającej. Producenci optycznych włókien i kapilar wysokiej jakości stosują testy wytrzymałości do całej produkcji [13]. Naprężenie we włóknie jest indukowane poprzez siły rozciągające, wyginające i skręcające. Przy naciągnięciu włókna naprężenie jest siłą przypadającą na jednostkę powierzchni poprzecznej szkła. Pokrycie polimerowe włókna nie przenosi naprężeń ale zabezpiecza włókno od narażeń mechanicznych i chemicznych, w tym wilgoci.



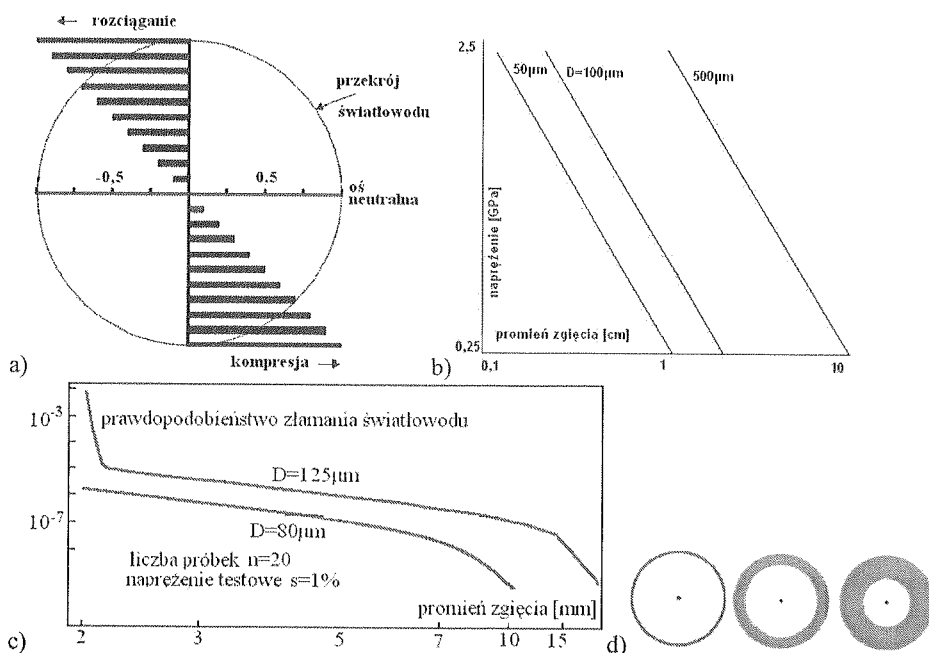
Rys. 13. Porównanie wytrzymałości światłowodów klasycznych krzemionkowych i ze szkła miękkich oraz kapilar. Pomiary własne. Dane od prawej: a) światłowod telekomunikacyjny komercyjny wielomodowy, krzemionkowy, rdzeń $50\mu\text{m}$, płaszcz $125\mu\text{m}$, pokrycie poliamidowe $150\mu\text{m}$; b) światłowod krzemionkowy kapilarny komercyjny, dane jak w a), średnica otworu $50\mu\text{m}$; c) światłowod krzemionkowy kapilarny, dane jak w a), średnica otworu $75\mu\text{m}$; d) światłowod ze szkła borokrzemionkowego miękkiego, dane jak w a), e) światłowod ze szkła barowego miękkiego, dane jak w a), f) światłowod kapilarny ze szkła borokrzemionkowego miękkiego, dane jak w c). Światłowody d), e) i f) wyciągano w KPO Politechniki Białostockiej

Fig. 13. Comparizon of tensile strengths of classical high silica and soft glass capillary optical fibres. Own measurements. Data from right: a) telecom grade optical fibre, multimode, core $50\mu\text{m}$, cladding $125\mu\text{m}$, jacket oliamide up to $150\mu\text{m}$; b) commercial capillary, data as in a), hole $50\mu\text{m}$; c) capillary optical fiber, data as in a), hole $75\mu\text{m}$; d) borosilicate glass optical fibre, data as ina, e) optical fibre from barium silica soft glass, data as in a), f) capillary optical fibre from borosilicate glass, data as in c), Fibres d), e) f) made in KPO Białystok Technical University

Napężenie wyginające można określić jako $\sigma = \varepsilon E = Er/(R + P_g + r)$ z następującymi parametrami: ε — naciąg, E — moduł Younga, r — promień światłowodu, P_g — grubość pokrycia światłowodu, R — promień wygięcia (krzywizny) światłowodu. Stąd $R = Er/\sigma - P_g - r$. Zależność pomiędzy promieniem wygięcia i napężeniem jest liniowa w skali podwójnie logarytmicznej. Profil napężenia w przekroju poprzecznym włókna jest w przybliżeniu $\sigma = (r_x/R)E$, gdzie r_x — jest odległością od osi neutralnej.

Podobny profil napężenie istnieje w światłowodzie kapilarnym. Różnica pomiędzy światłowodem pełnym i kapilarnym polega na istnieniu powierzchni wewnętrznej szkła, oprócz powierzchni zewnętrznej. Oś neutralna pod względem naprężeń znajduje się w dużej mierze nie w szkłe a w próżni, rys 14.d. W zależności od proporcji wymiarów

światłowodu kapilarnego (grubości ścianki) na wewnętrznej ścianie odkładane jest inne naprężenie, np. od 90% w przypadku cienkich ścianek do 50% w przypadku ścianek grubych. Oznacza to, że dla potencjalnego złamania spowodowanego defektem ścianki wewnętrznej, wymiar mikro-szczeliny musi być odpowiednio 1,1 lub 4 razy większy od mikro-szczeliny inicjalizującej na powierzchni zewnętrznej. Kapilary grubościennne, o tej samej średnicy otworu są mniej czułe na mikro-szczeliny na powierzchni wewnętrznej niż kapilary cienkościenne. W rzeczywistości, zależność pomiędzy promieniem zgięcia światłowodu kapilarnego i prawdopodobieństwem jego złamania jest bardziej skomplikowana gdyż musi uwzględnić naprężenia na powierzchni wewnętrznej otworu [12]. Rys. 14.c przedstawia wyniki obliczeń. Prawdopodobieństwo złamania kapilary maleje dla światłowodu o mniejszej średnicy zewnętrznej. Efekt grubości ścianki kapilary pokazano schematycznie na rys. 14.d) gdzie czułość na złamanie od mikro-szczeliny wewnętrznej jest redukowana odpowiednio 1,1, ok. 2 i ok. 4 razy.



Rys. 14. Naprężenia wyginające w światłowodzie. a) Rozkład naprężeń ściskających i rozciągających w przekroju poprzecznym wyginanego światłowodu, b) naprężenie maksymalne indukowane przez zgięcie, c) prawdopodobieństwo złamania światłowodu w wyniku wyginania, dla dwóch różnych średnic zewnętrznych włókna, d) proporcje wymiarowe kapilar o różnej grubości ścianek

Fig. 14. Bending stress in optical fibre. a) distribution of compressing and pulling stresses in the cross section of optical fibre, b) maximum stress induced by bending of optical fibre, c) probability of fracture of optical fibre as result of excessive bending for two different fibre diameters, d) dimensional proportions in capillary optical fibre of different wall thickness

5. PODSUMOWANIE

5.1. ZASADY UDOSTĘPNIANIA ŚWIATŁOWODÓW KSZTAŁTOWANYCH

Zespół technologiczny światłowodów kształtowanych ze szkła miękkich, reprezentowany przez autorów niniejszego opracowania grupuje laboratoria materiałowe, technologiczne, pomiarowe i aplikacyjne na terenie Politechniki Białostockiej i Politechniki Warszawskiej. Laboratoria współpracujące są także na terenie AGH. Zespół dysponuje opracowaniami światłowodami kształtowanymi, między innymi, następujących rodzajów: o złożonej refrakcji, o nietypowych kształtach rdzenia, wielordzeniowe oraz kapilarne.

Światłowody mogą być udostępniane zainteresowanym zespołom akademickim do badań w ilościach próbek laboratoryjnych bezpłatnie jedynie wtedy gdy ich poszukiwane rodzaje (wymiary, refrakcja, właściwości mechaniczne) są dostępne w składzie próbek laboratorium technologicznego. Specjalne żądania wykonania światłowodu kształtowanego, w tym kapilarnego, o parametrach dostosowanych do aplikacji kontrahenta jest wykonywane na zasadach komercyjnych w uzgodnionym terminie, na zasadzie zlecenia pracy technologicznej dla uczelni.

Laboratorium technologiczne światłowodów kształtowanych prowadzi ciągle badania nad nowymi rodzajami materiałów optycznych, włókien optycznych i podzespołów. Co pewien okres czasu uruchamiany jest specjalizowany proces technologiczny. Zainteresowane laboratoria zewnętrzne, po uzgodnieniu z laboratorium technologicznym, mogą uczestniczyć w tym procesie w celu uzyskania próbek optymalizowanych dla swoich potrzeb.

5.2. DALSZY PRACE PROWADZONE NAD ŚWIATŁOWODAMI KAPILARNYMI

Laboratorium technologiczne prowadzi dalsze prace nad światłowodami kapilarnymi. Związane jest to z dużym zainteresowaniem tymi światłowodami zespołów badawczych techniki światłowodowej. Zainteresowanie bierze się z dwóch zasadniczych powodów:

- kapilara optyczna jest elementarnym składnikiem światłowodu fotonicznego dziurawego;
- kapilary optyczne roszą nadzieje na zastosowania w dynamicznie rozwijającej się dziedzinie mikro i nanosystemów.

Prowadzona jest dalsza modernizacja systemów technologicznego nad rozszerzeniem możliwości wyciągania światłowodów przez laboratorium. Prowadzone są prace montażowe nad nową generacją pieca oporowego z precyzyjnym elementem grzejnym firmy Kanthal. Umożliwi to rozszerzenie zakresu temperatur wyciągania do ok. 1300°C i zapewni stabilniejszy rozkład temperatur wokół menisku wypływowego. Prowadzone są prace technologiczne nad precyzyjną kontrolą domieszkowania szkła miękkich pierwiastkami ziem rzadkich i metalami przejściowymi. Zastosowanie tych szkła do

budowy kapilar optycznych rokuje nadzieje na nieliniowe oddziaływania z zanikającą falą optyczną. Prowadzone są prace technologiczne nad precyzyjną kontrolą wyciągania osiowo symetrycznego układu kapilar w celu budowy światłowodu kapilarnego fotonicznego, w odróżnieniu od prezentowanych tutaj kapilar optycznych refrakcyjnych.

5.3. PARAMETRY WYTWARZANYCH ŚWIATŁOWODÓW KAPILARNYCH

Parametr	Rodzaj, wielkość
Rodzaje szkła	wieloskładnikowe sodowo-wapniowe, barowo-sodowe, ołowio-sodowe, borokrzemionkowe,
Technologia	tygłowa, preformowa, hybrydowa
Wymiary geometryczne	typoszerzeg wymiarowy
Średnica zewnętrzna [μm]	35 – 350
Średnica wewnętrzna [μm]	2 – 200
Stabilność wymiarów poprzecznych	mniej niż 3%
Średni stopień eliptyczności	mniej niż 1%
Pokrycie	poliamidowe, twarde,
Zakres przezroczystości [nm]	400 – 2500
Średni poziom przezroczystości	80%/m w zakresie widzialnym
Profil refrakcyjny	skokowy
Wytrzymałość mechaniczna	0,3-1,5 GPa
Parametr Weibulla	3 – 7
Średni łamiący promień zgięcia	ok. 5mm dla 50 μm do 30 mm dla 200 μm
Długość dostarczanych próbek	1 m, (specjalne do 5 m)

5.4. MOŻLIWOŚCI APLIKACJI ŚWIATŁOWODÓW KAPILARNYCH

Wraz z rozwojem techniki światłowodów kształtowanych rozszerzona została terminologia dotycząca także światłowodów kapilarnych. Nazywane są one również światłowodami z makrootworem, w odróżnieniu od światłowodów porowatych lub z mikrootworami. Często dodawana jest nazwa światłowody dziurawe na obydwie klasy włókien. Ogólnie można podzielić światłowody dziurawe z makrootworem na kilka grup: dziura-rdzeń; dziura obok rdzenia i rdzeń wokół dziury oraz rdzeń zawieszony w powietrzu. Początek tym technologiom dało zastosowanie kapilar — przetwarzanych na światłowody transmisyjne o ciekłych rdzeniach. Potencjalne możliwości zastosowań światłowodów kapilarnych są stosunkowo szerokie [8,9,10]. Związane jest to z następującymi grupami czynników i zjawisk:

- możliwość transmisji dużej mocy optycznej w powietrznym rdzeniu,
- oddziaływanie bliskiego otworu na rdzeń światłowodu,
- możliwość wypełnienia powietrznego rdzenia substancją mierzoną;
- możliwość tworzenia długiej drogi oddziaływania pomiędzy falą zanikającą rdzenia z otworem,
- możliwość prowadzenia w otworze fali powierzchniowej typu WGM (whispering gallery modes),

— możliwości wykorzystania zjawiska modulacji napięcia powierzchniowego cieczy wewnątrz kapilary,

Ostatnim znacznym osiągnięciem zastosowań światłowodów kapilarnych jest koherentna transmisja fali materialnej de Broglia, wielomodowej i jednomodowej. Polega ona na uporządkowanej transmisji, w polu optycznym światłowodu kapilarnego, strumienia pojedynczych atomów. Światłowodowe stożki kapilarne zastosowano do budowy optycznej wersji mikroskopu sił atomowych (AFM). Światłowodowe stożki kapilarne stosowane są do próbkowania badanych powierzchni atomami helu oraz do układania pojedynczych warstw atomów.

Światłowodowy kapilarny są zupełnie odmienne we właściwościach i zastosowaniach od klasycznych światłowodów. Ich podstawową cechą aplikacyjną jest możliwość zbliżenia się z oddziaływaniem zewnętrznym (pozostając wewnątrz światłowodu) do transmisyjnego optycznego pola rdzeniowego. Poprzez bliskość takich oddziaływań, światłowodowy kapilarny mogą potencjalnie wykazywać znacznie większe wrażliwości od włókien o konstrukcjach klasycznych. Najszerszym przykładem potencjalnej aplikacji jest zintegrowany spektrometr światłowodowy, działający na fali zanikającej, gdzie ośrodek podlegający badaniom znajduje się wewnątrz włókna. Takie światłowodowy mogą stać się elementami mikrosystemów.

Badania aplikacyjne nad światłowodami kapilarnymi wykorzystują najczęściej włókna optyczne o standaryzowanych wymiarach zewnętrznych wynoszących 125 lub 150 um. Wynika to z faktu konieczności wykorzystania do badań adaptowanego sprzętu telekomunikacyjnego. Najczęściej spotykane wymiary otworów (pojedynczego lub podwójnego) wewnątrz światłowodu wynoszą kilkadziesiąt um. Ze światłowodów kapilarnych lub o konstrukcji analogicznej do kapilarnych (np. side hole) budowano szereg urządzeń:

- czujniki temperatury i naprężeń,
- elementy o silnie wyrażanych niektórych zjawiskach optycznych, np. elastooptycznych,
- polaryzatory światłowodowe,
- czujniki chemiczne, optrody, amplitudowe, spektrometryczne,
- ostatnio, elementy mikrosystemów.

Podstawowe parametry aplikacyjne światłowodów kapilarnych, istotne dla użytkownika i projektanta czujników i elementów funkcjonalnych z tych światłowodów to: materiał światłowodu, właściwości refrakcyjne, parametry geometryczne jak — wymiary zewnętrzne oraz proporcje wymiarów pomiędzy otworem i rdzeniem optycznym. Te parametry są dostarczane użytkownikom przez producentów włókien, w ramach danych technicznych, łącznie z próbkami włókien. W zależności od materiału światłowodu, włókno kapilarne może być uczulane na poszczególne oddziaływania zewnętrzne. W zależności od poziomu refrakcji, włókno może być stosowane w różnych środowiskach refrakcyjnych. W zależności od relacji wymiarów, włókno może być stosowane jako czujnik amplitudowy lub fazowy.

6. PODZIĘKOWANIA

Niniejsza praca była częściowo finansowana w ramach programu PBZ-MIN-009/T11/2003 pt. Elementy i moduły optoelektroniczne do zastosowań w medycynie, przemyśle, ochronie środowiska i technice wojskowej.

7. BIBLIOGRAFIA

1. Ch.Y.H. Tsao, D.N. Payne, W.A. Gambling: *Modal characteristics of three-layered optical fiber waveguides: a modified approach*, J.Opt.Soc.Am.A., vol. 6, no. 4, April 1989, pp. 555-563.
2. M. Hautakorpi, M. Kaivola: *Modal analysis of M-type-dielectric-profile optical fibers in the weakly guiding approximation*, J.Opt.Soc.Am.A., vol. 22, no. 6, June 2005, pp. 1163-1169.
3. M. Ibanescu, S.G. Johnson, M. Soljacic, J.D. Joannopoulos, Y. Fink, O. Weisberg, T.D. Engeness, S.A. Jacobs, M. Skorobogatiy: *Analysis of mode structure in hollow dielectric waveguide fibers*, Physical Review, E 67, 046608 (2003), pp. 1-8.
4. A.D. Fitt, K. Furusawa, T.M. Monro, C.P. Please: *Modelling the fabrication of hollow fibers: capillary drawing*, Journ. Lightwave Technol., vol. 19, no. 12, December 2001, pp. 1924-1930; (i literatura tam zamieszczona).
5. Ch.D. Rabi, J.A. Harrington: *Mechanical properties of hollow glass waveguides*, Opt.Eng., 38(9), pp. 1490-1499 (September 1999).
6. R. Romaniuk, J. Dorosz: *Kontrola geometrii światłowodów kapilarnych*, Elektronika, nr 4, 2006; (i literatura tam zamieszczona).
7. R. Romaniuk, J. Dorosz: *Właściwości mechaniczne światłowodów kapilarnych*, Elektronika nr 3, 2006; (i literatura tam zamieszczona).
8. R. Romaniuk, J. Dorosz: *Transmisja koherentnej fali deBroglia w światłowodzie kapilarnym*, Elektronika, nr 3, 2006; (i literatura tam zamieszczona).
9. R. Romaniuk: *Zastosowania światłowodów kapilarnych*, Elektronika nr. 4, 2004.
10. R. Romaniuk: *Światłowodowy kapilarne w telekomunikacji*, Elektronika, nr 6, 2006; (i literatura tam zamieszczona).

R. ROMANIUK, J. DOROSZ

FABRICATION AND CHARACTERIZATION OF CAPILLARY OPTICAL FIBRES

Summary

The paper describes basic properties of capillary optical fibres. The analysis of propagation properties were shown. The capillaries were fabricated of soft glasses by crucible and perform methods. The used glasses were borosilicates and calcium-sodium silicates. The following characteristics were measured: geometrical, optical and mechanical. The calculations were compared with measurements.

Keywords: tailored optical fibres, optical capillaries, optical fiber technology, measurements of optical fibres

Projektowanie topografii systemów VLSI

Cz. 1. Style i etapy projektowania, rozmieszczanie modułów

ZBIGNIEW NAGÓRNY, ANDRZEJ KOS

*Zbigniew Nagórny,
Studium Generale Sandomiriense,
Wyższa Szkoła Humanistyczno-Przyrodnicza w Sandomierzu,
ul. Krakowska 26, 27-600 Sandomierz,
e-mail: nagorny@interia.pl,*

*Andrzej Kos,
Akademia Górniczo-Hutnicza w Krakowie, Katedra Elektroniki,
al. Mickiewicza 30, 30-059 Kraków,
e-mail: kos@agh.edu.pl*

*Otrzymano 2005.10.06
Autoryzowano 2006.02.20*

Projektowanie układów VLSI wymaga stosowania systemów projektowania wspomaganych komputerowo. Niniejsza praca jest pierwszą częścią przeglądu metod rozmieszczania modułów, stosowanych podczas projektowania topografii układów VLSI. Opisano różne style topografii oraz przykłady układów dla poszczególnych stylów. Następnie, przedstawiono etapy projektowania topografii: podział, planowanie układu, rozmieszczanie, trasowanie połączeń oraz weryfikacja. Planowanie układu zostało szczegółowo omówione, ze względu na podobieństwa łączące ten etap z rozmieszczaniem. Przedstawiono problem rozmieszczania modułów. Omówiono sposoby estymacji długości połączeń. Opisano metody minimalizacji opóźnień w układzie. Przedstawiono stosowane metody rozmieszczania modułów.

Słowa kluczowe: VLSI, projektowanie topografii układu, układ scalony, projektowanie wspomagane komputerowo, układ Standard Cell, matryca bramkowa GA, układ Sea-of-Gates, układ FPGA, podział, planowanie układu, rozmieszczanie modułów, trasowanie połączeń, estymacja długości połączeń, metoda połowy obwodu, graf pełny, minimalizacja opóźnień, metody konstrukcyjne, algorytm min-cut, algorytm zamiany parami, symulowane wyżarzanie, algorytm genetyczny, metody analityczne, sieć neuronowa

1. WPROWADZENIE

Projektowanie topografii cyfrowego układu VLSI (ang. physical design process) jest bardzo trudnym problemem optymalizacyjnym. Rezultatem tego procesu jest wiedza prowadząca do fizycznej realizacji projektowanego układu. Projekt topografii (ang. layout) określa przede wszystkim obszary dyfuzji oraz metalizacji w układzie scalonym. Ogólnie można powiedzieć, że celem optymalizacji topografii układu VLSI jest minimalizacja jego powierzchni oraz zapewnienie możliwie najlepszych właściwości funkcjonalnych układu. Mniejsza powierzchnia układu pozwala uzyskać większy uzysk podczas produkcji układu, ponieważ umożliwia wykonanie większej liczby układów na jednostkowej powierzchni podłoża oraz powoduje zmniejszenie prawdopodobieństwa wystąpienia defektu podłoża w pojedynczym układzie. Podczas projektowania topografii należy również wziąć pod uwagę minimalizację opóźnień w krytycznych częściach układu, minimalizację traconej energii, sprzężeń między sygnałami oraz wiele innych kryteriów [1 – 4, 38].

Ze względu na rozmiary i złożoność obecnych układów VLSI, projektowanie topografii układów musi odbywać się z użyciem systemów projektowania wspomaganych komputerowo CAD (ang. computer aided design). Systemy te umożliwiają znaczne skrócenie czasu niezbędnego do zaprojektowania układu, poprawienie jego jakości i niezawodności oraz zmniejszenie kosztów jego produkcji [3]. Współczesne systemy projektowania zawierają wiele narzędzi, których zadaniem jest usprawnienie procesu projektowania. Wśród tych narzędzi należy wymienić edytory topografii układu, narzędzia sprawdzające zgodność topografii z regułami projektowania i schematem elektrycznym układu, narzędzia generujące rzeczywisty model układu z uwzględnieniem pojemności i rezystancji pasożytniczych, pojemności, indukcyjności i rezystancji połączeń oraz inne narzędzia [1, 3, 5].

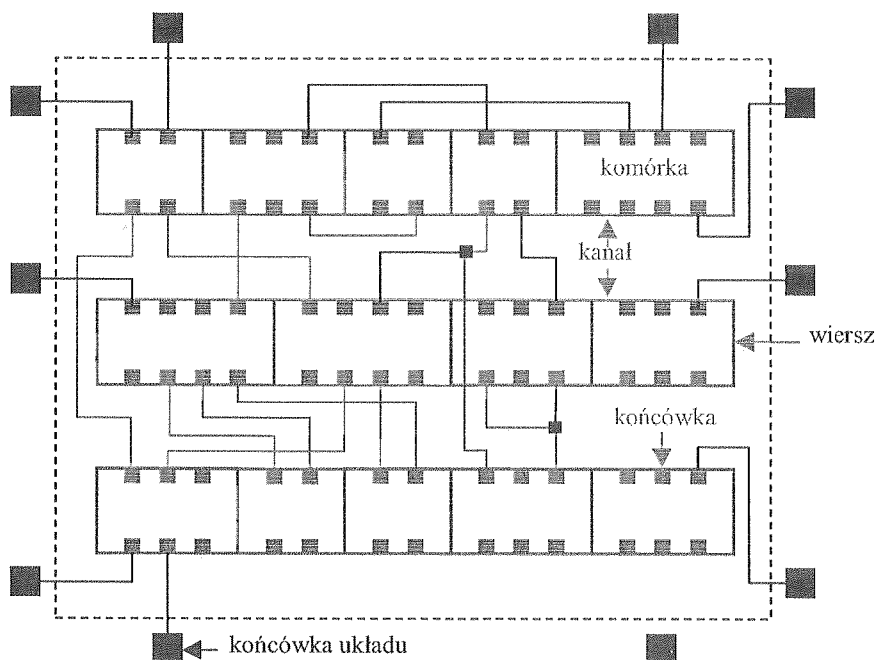
Niniejsza praca jest pierwszą częścią przeglądu metod rozmieszczania modułów, stosowanych podczas projektowania topografii układów VLSI. Stosowane style topografii układów przedstawione są w rozdziale 2. Rozdział 3 zawiera opis poszczególnych etapów projektowania topografii układu. Problem rozmieszczania modułów w projektowaniu topografii układów VLSI jest przedstawiony w rozdziale 4. Podsumowanie pierwszej części przeglądu metod rozmieszczania zamieszczono w rozdziale 5. Bibliografia znajduje się w rozdziale 6.

2. STYLE TOPOGRAFII UKŁADU VLSI

Systemy projektowania topografii układów VLSI muszą uwzględniać styl (strukturę) topografii układu (ang. design style). Topografia układu VLSI może zostać zaprojektowana w jednym z trzech stylów: topografia o strukturze wierszowej — topografia wierszowa (ang. row based), topografia o strukturze swobodnej — topografia swobodna (ang. building block) oraz topografia mieszana (ang. mixed-size, mixed block design). W przypadku topografii wierszowej układ składa się z wierszy, w których umieszczono

ne są elementarne komórki logiczne. Topografię wierszową posiadają: układy Standard Cell, matryce bramkowe GA (ang. Gate Array), układy Sea-of-Gates SOG, układy FPGA (ang. Field Programmable Gate Array).

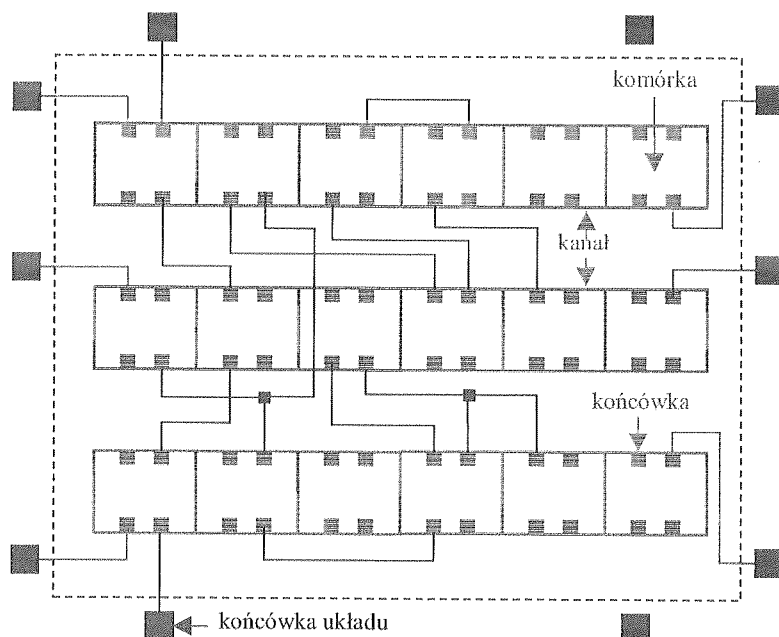
Topografia układów Standard Cell jest przedstawiona na rysunku 1. Układ składa się z wierszy elementarnych komórek logicznych — komórek standardowych, które zostały wcześniej zaprojektowane przez producenta układów. Projektant danego układu realizuje jego funkcję logiczną, używając tylko komórek standardowych zebranych w bibliotece. Komórki standardowe mają przeważnie taką samą wysokość, natomiast szerokość komórek może być różna. Połączenia między komórkami prowadzone są w kanałach, które powinny posiadać wysokość niezbędną do poprowadzenia wszystkich połączeń. Jeżeli połączenie musi przeciąć wiersz komórek, wówczas może być ono poprowadzone pionowo w specjalnej części jednej z komórek lub z wykorzystaniem specjalnej komórki, przeznaczonej wyłącznie do prowadzenia takich połączeń (ang. feedthrough). Układy Standard Cell wymagają wykonania wszystkich fotomasek przez producenta układu. Dzięki temu możliwe jest dostosowanie rozmiarów tranzystorów do potrzeb konkretnego układu, czynność ta odbywa się na poziomie biblioteki komórek. Wadą układów Standard Cell jest wysoki koszt oraz długi okres produkcji układu, w porównaniu do matryc bramkowych GA, układów Sea-of-Gates oraz układów programowalnych [1, 3, 6–9, 36].



Rys. 1. Topografia układu Standard Cell

Fig. 1. Standard Cell layout

Topografię matryc bramkowych GA przedstawia rysunek 2. Układy te również składają się z wierszy elementarnych komórek logicznych, które zostały wcześniej zaprojektowane przez producenta układów. Wszystkie komórki w układzie posiadają taką samą topografię. Połączenia prowadzone są w kanałach o ustalonej wysokości. Z tego względu produkowane są układy o różnej wysokości kanału. Produkcja układów może być masowa, co wpływa na zmniejszenie kosztu realizowanego układu. Projektant konkretnego układu określa jedynie topografię połączeń między poszczególnymi komórkami, co umożliwia skrócenie czasu jego realizacji, ponieważ producent wykonuje specjalnie dla tego układu jedynie fotomaski metalizacji. Wadą matryc bramkowych GA jest mniejsze wykorzystanie powierzchni układu oraz gorsze właściwości funkcjonalne układu, ze względu na brak możliwości dostosowania poszczególnych tranzystorów do potrzeb danego układu [1, 3, 6–7, 9].

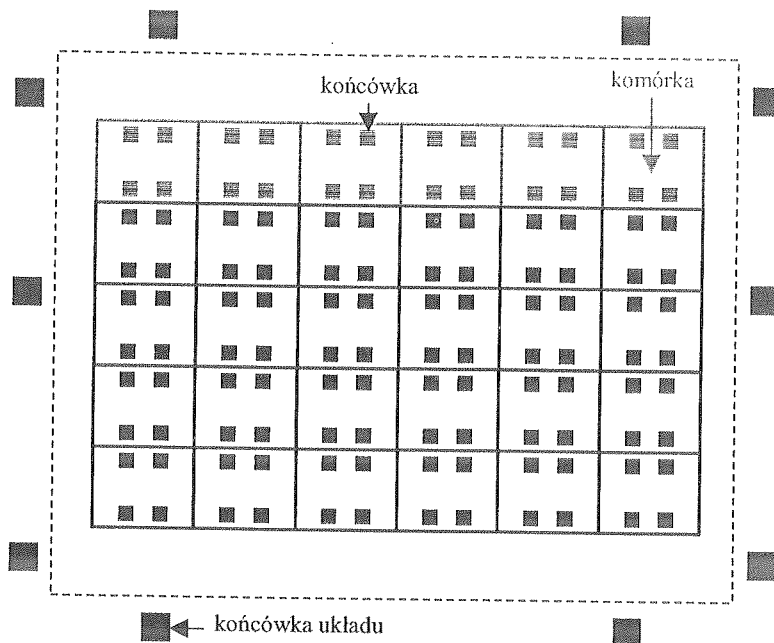


Rys. 2. Topografia matrycy bramkowej GA

Fig. 2. Gate Array layout

Rysunek 3 przedstawia topografię układów Sea-of-Gates. Układy te są podobne do matryc bramkowych GA, ale nie posiadają wyznaczonych kanałów. Cała powierzchnia układu składa się z elementarnych komórek logicznych, które posiadają taką samą topografię, zaprojektowaną przez producenta układów. Podczas projektowania konkretnego układu, określona liczba wierszy komórek jest rezerwowana do poprowadzenia połączeń. Komórki te nie realizują wówczas żadnych funkcji logicznych i stanowią kanały. Układy Sea-of-Gates charakteryzują się lepszym wykorzystaniem powierzchni układu w stosunku do matryc bramkowych GA [1, 3]. Układy Standard Cell, matryce

bramkowe GA oraz układy Sea-of-Gates określamy mianem układów typu semi custom (ang.), ze względu na dwuetapowy cykl projektowania układu.



Rys. 3. Topografia układu Sea-of-Gates

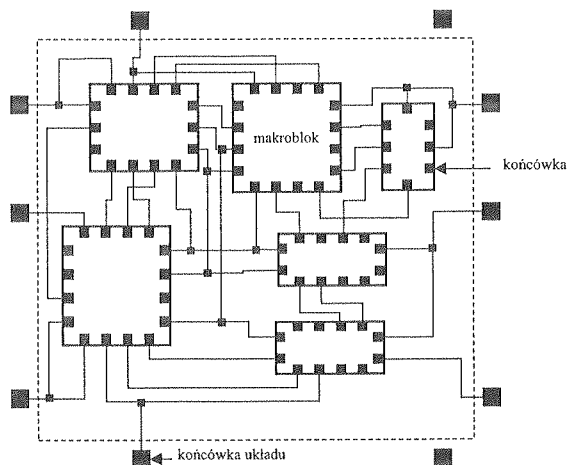
Fig. 3. Sea-of-Gates layout

Przykład topografii swobodnej jest przedstawiony na rysunku 4. Części składowe układu — makrobloki (ang. macro block), mogą posiadać dowolne położenie i rozmiar, co umożliwia większe wykorzystanie powierzchni układu w stosunku do układów o topografii wierszowej. Przestrzenie między makroblokami służą do prowadzenia połączeń. Układy te są projektowane od podstaw, co umożliwia również osiągnięcie lepszych właściwości funkcjonalnych układu. Z tego względu określamy je mianem układów typu full-custom (ang.). Układy te wymagają wykonania wszystkich fotomasek przez producenta układu [3, 6–7, 9].

Rysunek 5 przedstawia przykład topografii mieszanej. Układ składa się z dwóch części, posiadających różną topografię: wierszową (układ Standard Cell) oraz swobodną (makrobloki). Połączenie tych topografii w jednym układzie jest możliwe, ponieważ układy Standard Cell oraz topografia swobodna wymagają wykonania wszystkich fotomasek przez producenta układu [1, 8–9, 36].

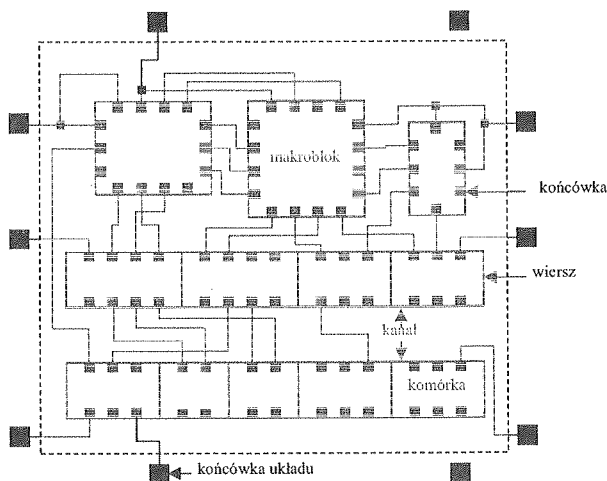
Układy FPGA są produkowane we wszystkich trzech stylach topografii. Rysunek 6 przedstawia typową topografię układu FPGA. Układ składa się z prostokątnej macierzy elementarnych komórek logicznych. Realizacja projektu konkretnego układu odbywa się programowo, bez udziału producenta. Programowanie układu odbywa się w wyni-

ku programowania elementarnych komórek logicznych CLB (ang. Configurable Logic Block) oraz programowania połączeń między nimi. Połączenia między komórkami logicznymi prowadzone są z użyciem poziomych oraz pionowych kanałów. Programowanie połączeń między komórkami zapewniają programowalne przełączniki, znajdujące się w rogach między czterema komórkami logicznymi. Połączenia z końcówkami całego układu są realizowane za pomocą programowalnych komórek wejścia/wyjścia I/O (ang. Input-Output Block) [1, 10, 45–46].



Rys. 4. Topografia swobodna

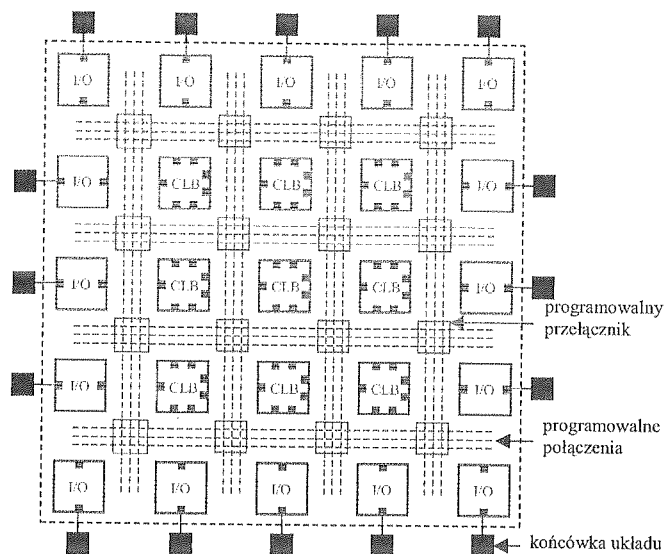
Fig. 4. Layout of a building block design style



Rys. 5. Topografia mieszana

Fig. 5. Layout of a mixed design style

ogic
tami
mo-
jdu-
kami
jsia



Rys. 6. Topografia typowego układu FPGA

Fig. 6. Layout of a typical FPGA circuit

Metody rozmieszczania stosowane podczas projektowania układów VLSI muszą uwzględniać styl topografii oraz typ projektowanego układu. Niektóre metody zostały opracowane dla konkretnego typu układów. Inne mogą być stosowane dla różnych układów, po odpowiednich modyfikacjach.

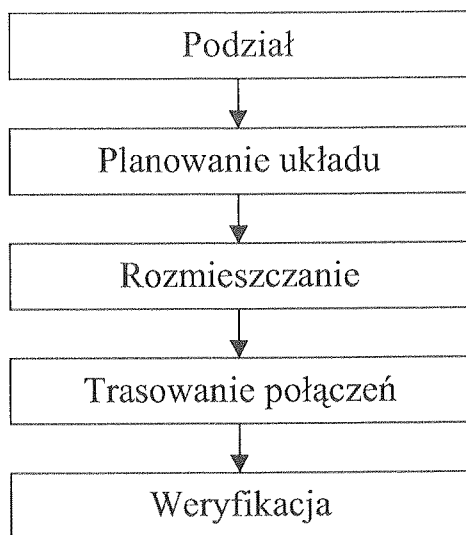
3. ETAPY PROJEKTOWANIA TOPOGRAFII UKŁADU

Wejściowymi danymi dla procesu projektowania topografii układu VLSI jest lista połączeń części składowych układu (ang. netlist). Częściami składowymi układu mogą być pojedyncze tranzystory, bramki logiczne, komórki standardowe, w zależności od poziomu hierarchii, na jakim rozpatrujemy projektowany układ.

W obecnych systemach automatycznego projektowania topografii układów VLSI, proces projektowania jest podzielony na kilka etapów, ponieważ nie jest możliwe rozwiązanie tego problemu w jednym etapie. Proces projektowania topografii układu składa się zatem z następujących etapów: podziału (ang. partitioning), planowania układu (ang. floorplanning, chip planning), rozmieszczania (ang. placement), trasowania połączeń (ang. routing) oraz weryfikacji (ang. verification) [1, 3–7]. Rysunek 7 przedstawia etapy projektowania topografii układu VLSI.

Podział może być pierwszym etapem projektowania topografii układu, gdy rozmiar układu nie pozwala na wykonanie go w jednym układzie scalonym. Podział umożliwia zastąpienie całego układu pewną liczbą układów scalonych. Polega na określeniu,

które części składowe całego układu powinny znaleźć się w poszczególnych układach scalonych. Wykorzystywana jest tutaj metoda „dziel i rządź” (ang. „divide and conquer”), która jest bardzo często stosowana w procesie projektowania topografii. Metoda ta polega na podziale problemu na mniejsze podproblemy. Dzięki temu możliwe jest otrzymanie rozwiązania również dla bardzo dużych układów, dla których niemożliwe jest uzyskanie rozwiązania w sposób bezpośredni [1, 3–5].

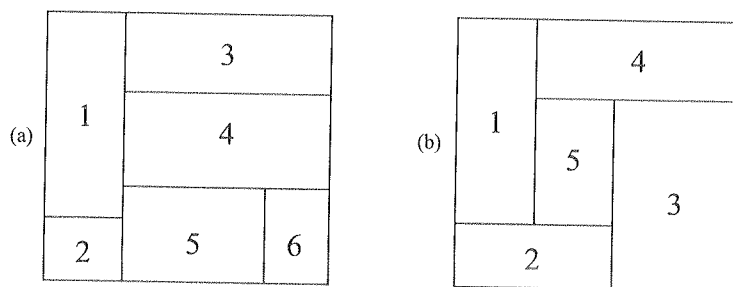


Rys. 7. Etapy projektowania topografii układu VLSI

Fig. 7. Physical design flow

Kolejnym etapem jest planowanie układu, podczas którego następuje wstępne określenie topografii układu scalonego, w wyniku rozmieszczenia poszczególnych bloków funkcjonalnych układu. Podczas tego etapu określany jest sposób prowadzenia zasilania, sygnału taktującego oraz masy w układzie. Określane jest również położenie poszczególnych końcówek całego układu (ang. pad). Podczas tego etapu rozpatrywane są bloki, dla których rozmiary oraz położenie końcówek jest określone (ang. hard block) oraz bloki, które mają jedynie określoną powierzchnię, jaką będą zajmować w układzie (ang. soft block). Bloki typu soft nie posiadają określonego położenia końcówek, natomiast stosunek wysokości i szerokości tego bloku (ang. aspect ratio) może zmieniać się w zadanym przedziale [1, 3, 5, 7, 11, 47]. Jeżeli możliwy jest podział układu oraz jego części liniami poziomymi i pionowymi tak, aby każdy blok stanowił w końcu jedną część, wówczas plan topografii układu jest typu slicing (ang.). W przeciwnym razie jest typu non-slicing (ang.) [1, 7–8, 48]. Rysunek 8 przedstawia plan topografii układu typu slicing oraz non-slicing. Ze względu na podobieństwa łączące planowanie układu oraz rozmieszczanie, metody stosowane podczas tego etapu zostaną dokładniej opisane.

Ze względu na dużą złożoność obliczeniową, planowanie układu odbywa się zwykle w dwóch krokach [7, 11, 47]. Podczas pierwszego kroku następuje określenie wzajemnego położenia bloków tak, aby zminimalizować całkowitą długość połączeń w układzie. Czynność ta odbywa się z wykorzystaniem teorii grafów [9, 11–13, 18]. W drugim kroku określone jest położenie poszczególnych bloków oraz rozmiary bloków typu soft tak, aby zminimalizować całkowitą powierzchnię układu, zachowując wcześniej określone wzajemne położenie bloków. W tym celu stosowane są metody wykorzystujące algorytm Ottena i Stockmeyera [14–16], algorytm podziału i ograniczenia (ang. branch and bound) [14–15], teorię grafów [13, 16–18], podział planu topografii układu [19], programowanie liniowe (ang. linear programming LP) [9, 11–12, 47].



Rys. 8. Plan topografii układu: a) typu slicing, b) typu non-slicing

Fig. 8. Floorplans: a) slicing, b) non-slicing

Istnieją również metody, które pozwalają na minimalizację całkowitej powierzchni układu oraz długości połączeń w jednym kroku. Metody te najczęściej wykorzystują symulowane wyżarzanie, które wymaga odpowiedniego sposobu kodowania rozwiązania. Bardzo istotne jest, aby zmiany rozwiązania dokonywane podczas symulowanego wyżarzania nie powodowały powstawania rozwiązań, które są fizycznie niemożliwe do wykonania [7, 11, 20]. W przypadku planu topografii typu slicing, do kodowania rozwiązań stosowana jest Odwrotna Notacja Polska (ang. reverse Polish notation) [7]. Dla planu topografii typu non-slicing bardzo często stosowane jest kodowanie z użyciem pary nazw bloków, określających ich kolejność (ang. sequence-pair) [20–22]. Plan topografii typu non-slicing zapewnia lepsze upakowanie bloków w układzie, lecz wymaga dodatkowych nakładów podczas trasowania połączeń [8]. Obecnie badania są skupione głównie na znalezieniu efektywnego sposobu kodowania rozwiązań [23–24, 48]. Odmienny sposób podejścia do problemu planowania topografii układu, w którym rozmiary planowanego układu są ściśle określone, można znaleźć w pracy [25].

Kolejnym etapem jest rozmieszczanie. Podczas tego etapu określone jest położenie wszystkich części składowych układu w poszczególnych blokach układu. W dalszej części przeglądu części składowe układu będą nazywane modułami. Najczęściej stosowanymi kryteriami podczas rozmieszczania modułów jest minimalizacja całkowitej

tej estymowanej (przewidywanej) długości połączeń, minimalizacja opóźnień w krytycznych częściach układu oraz minimalizacja skupienia połączeń [1, 3–7, 9, 38–39]. Rozmieszczanie zwykle składa się z dwóch kroków: rozmieszczania globalnego (ang. global placement) i rozmieszczania szczegółowego (ang. detailed placement). Podczas rozmieszczania globalnego moduły są prawie równomiernie rozmieszczane na podłożu układu, zgodnie z wybranymi kryteriami rozmieszczania. Rozmieszczanie szczegółowe ma na celu wyznaczenie ostatecznego położenia modułów, zgodnie z wybranym stylem topografii układu, usuwane jest zachodzenie modułów na siebie. Podczas rozmieszczania szczegółowego może być również wykonywana korekcja rozmieszczenia dla małych grup modułów, według wybranych kryteriów rozmieszczania. Niniejszy przegląd jest poświęcony metodom rozmieszczania, dlatego metody te zostaną szczegółowo opisane w kolejnych rozdziałach. Planowanie układu i rozmieszczanie łączy wiele podobieństw, z tego względu istnieje wiele metod, które z powodzeniem mogą być stosowane podczas obydwu etapów [3, 5].

Po rozmieszczeniu modułów możliwe jest wyznaczenie połączeń między nimi. Czynność ta odbywa się podczas etapu trasowania połączeń. Etap ten składa się z dwóch kroków: trasowania globalnego (ang. global routing) i trasowania szczegółowego (ang. detailed routing). Podczas trasowania globalnego połączenia nie są wyznaczane, a jedynie planowane. Planowanie połączeń może odbywać się z wykorzystaniem grafu, którego wierzchołkami są punkty przecięcia kanałów, w których mają być prowadzone połączenia. Krawędziami grafu są kanały. Wagi krawędzi powinny odzwierciedlać długość połączeń oraz pojemność kanałów, powinny być aktualizowane po wyznaczeniu poszczególnych połączeń. W celu wyznaczenia danego połączenia, w grafie wyznaczana jest ścieżka o najmniejszym koszcie [1, 3–5, 8]. Do trasowania globalnego może być zastosowany również algorytm symulowanego wyżarzania [6]. Po zaplanowaniu połączeń, następuje fizyczne wyznaczenie połączeń podczas trasowania szczegółowego. W tym kroku połączenia są prowadzone w kanałach, które zostały wcześniej przydzielone poszczególnym połączeniom podczas trasowania globalnego. Istnieje cały szereg metod trasowania szczegółowego [1, 3–5, 7–8, 10, 27–28].

Po wyznaczeniu topografii układu następuje weryfikacja rozwiązania. Obejmuje ona sprawdzenie zgodności projektu topografii z regułami projektowania i schematem elektrycznym układu. Może polegać również na określeniu rzeczywistego modelu układu. Znajomość pojemności i rezystancji pasożytniczych, pojemności, indukcyjności i rezystancji połączeń oraz parametrów tranzystorów, umożliwia dokładną symulację układu [1, 3–5, 26]. Przed etapem weryfikacji topografii układu stosowana jest również kompresja (zagęszczanie) topografii (ang. compaction). Celem kompresji topografii jest zmniejszenie rozmiarów układu poprzez eliminację niepotrzebnych przestrzeni w układzie, przy zachowaniu reguł projektowania oraz topologii układu. Kompresja może być stosowana również w przypadku zmiany technologii wykonania układu [3, 5].

Czasową złożoność obliczeniową algorytmu mierzy się liczbą elementarnych operacji powtarzanych w algorytmie. Stosowane jest pojęcie złożoności asymptotycznej. Niech n oznacza rozmiar problemu. Algorytm posiada złożoność $O(f(n))$, jeżeli istnieją

liczby dodatnie c i N takie, że liczba elementarnych operacji jest mniejsza lub równa niż $cf(n)$ dla wszystkich $n \geq N$. Wynika z tego, że dla dostatecznie dużych n liczba elementarnych operacji rośnie nie szybciej niż funkcja f . Większość etapów projektowania topografii układu (z wyjątkiem weryfikacji) zaliczanych jest do tzw. problemów NP-trudnych. Dla tych problemów nie znaleziono algorytmu wyznaczenia optymalnego rozwiązania, którego czasowa złożoność obliczeniowa jest wielomianem dowolnego, skończonego stopnia względem rozmiaru problemu. Z tego powodu do rozwiązania problemu stosowane są metody, które są aproksymacją optymalnego rozwiązania [1, 3-7, 9, 14, 16-20, 24, 30].

4. PROBLEM ROZMIESZCZANIA MODUŁÓW

W niniejszym przeglądzie części składowe układów będą nazywane modułami. W ten sposób będą określane zarówno komórki standardowe w układach Standard Cell, jak również elementarne komórki logiczne w pozostałych układach o topografii wierszowej. Części składowe układów o topografii swobodnej również będą nazywane modułami.

4.1. POSTAWIENIE PROBLEMU

Rozmieszczanie modułów w układzie VLSI jest bardzo ważnym etapem projektowania jego topografii. Sposób rozmieszczenia modułów decyduje o powierzchni układu oraz jego właściwościach funkcjonalnych. Podczas następnych etapów projektowania topografii nie można naprawić błędów popełnionych podczas rozmieszczania [32]. Danymi wejściowymi dla metod rozmieszczania są kształty, rozmiary i położenie końcówek poszczególnych modułów układu, położenie końcówek całego układu oraz lista połączeń między modułami i końcówkami całego układu. Rozwiązanie problemu rozmieszczania polega na znalezieniu współrzędnych prostokątnych, określających położenie modułów w układzie tak, aby zapewnić spełnienie celów stawianych procesowi rozmieszczania. Metody rozmieszczania muszą uwzględniać styl topografii układu, ponadto moduły nie mogą zachodzić na siebie oraz w całości muszą być położone w obszarze układu. W układach o topografii wierszowej często zakłada się, że wszystkie końcówki modułów są położone w środku geometrycznym modułów, ponieważ rozmiary modułów są niewielkie w porównaniu do rozmiarów podłoża układu. W układach o topografii swobodnej położenie końcówek modułów (makrobloków) musi być uwzględnione, ponieważ poszczególne moduły mogą znacznie różnić się od siebie rozmiarami i zajmować znaczną część podłoża układu. Uwzględniana jest również orientacja modułu (makrobloku). Każdy moduł może posiadać osiem różnych orientacji, które są rezultatem obrotu oraz odbicia względem jego osi symetrii [38, 49].

Podstawowymi celami rozmieszczania jest umożliwienie poprowadzenia wszystkich połączeń, minimalizacja powierzchni układu, minimalizacja opóźnień w krytycznych częściach układu, minimalizacja sprzężeń między sygnałami, minimalizacja wy-

dzielanej energii w układzie, równomierny rozkład temperatury układu oraz wiele innych [1–5, 9, 30, 38–39]. Cele te trudno jest formalnie zawrzeć w algorytmach rozmieszczania, z tego względu metody rozmieszczania najczęściej stosują własne kryteria. Najczęściej stosowanymi celami metod rozmieszczania jest minimalizacja całkowitej estymowanej długości połączeń, minimalizacja skupienia połączeń, minimalizacja opóźnień w krytycznych częściach układu oraz minimalizacja powierzchni układu — w przypadku topografii swobodnej. Cele rozmieszczania są zwykle określone przez tzw. funkcję kosztu. Rozwiązanie problemu polega na znalezieniu rozmieszczenia, które posiada minimalny koszt [1, 3–10, 12, 49].

Minimalizacja skupienia połączeń pośrednio wpływa na minimalizację powierzchni układu, ze względu na mniejsze rozmiary kanałów do prowadzenia połączeń oraz pośrednio przyczynia się do minimalizacji długości połączeń, ponieważ silnie połączone ze sobą moduły są rozmieszczane obok siebie [1, 5, 8–9, 31, 33–35, 37]. Minimalizacja całkowitej estymowanej długości połączeń jest bardzo często stosowanym celem metod rozmieszczania. Minimalizacja długości połączeń pośrednio przyczynia się do minimalizacji całkowitej powierzchni układu, ponieważ połączenia zajmują znaczną część jego powierzchni, nawet 50–70% całkowitej powierzchni [8–9, 29–30, 31–32]. Pośrednio przyczynia się do minimalizacji skupienia połączeń i umożliwia wykonanie wszystkich połączeń, w przypadku ograniczonej pojemności kanałów przeznaczonych do prowadzenia połączeń [4, 10, 30]. Powoduje zmniejszenie pojemności i rezystancji pasożytniczych, wnoszonych do układu przez połączenia oraz pojemności, indukcyjności i rezystancji połączeń. Minimalizacja długości połączeń umożliwia zatem również zmniejszenie czasu propagacji sygnałów oraz ilości ciepła wydzielanego w układzie. Minimalizacja opóźnień w krytycznych częściach układu zapewnia znaczne zmniejszenie czasu propagacji sygnałów w układzie, co umożliwia zastosowanie sygnału taktującego o większej częstotliwości [1, 3, 9, 40–44].

4.2. MINIMALIZACJA DŁUGOŚCI POŁĄCZEŃ

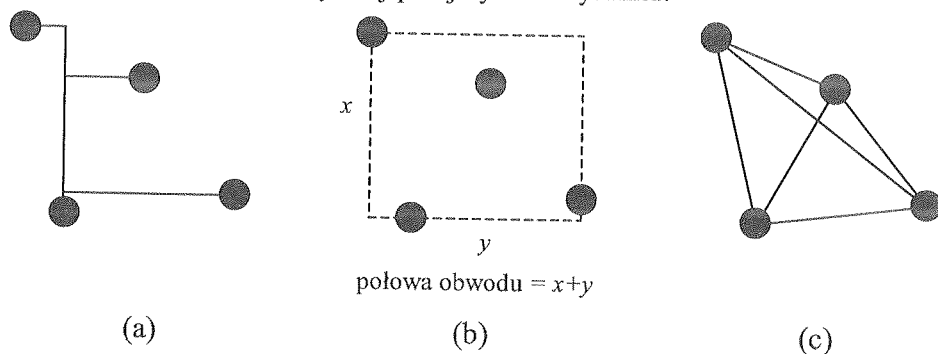
Połączenia w układach VLSI prowadzone są zwykle wyłącznie w kierunku poziomym oraz pionowym, dlatego do określenia długości połączeń stosowana jest najczęściej metryka Manhattan, według wzoru

$$d_{ij} = |x_i - x_j| + |y_i - y_j| \quad (1)$$

gdzie: d_{ij} — długość połączenia między punktami o współrzędnych (x_i, y_i) oraz (x_j, y_j) . Wykorzystywana jest również metryka Euklidesa oraz kwadrat długości euklidesowej [1, 3, 9, 28, 30].

Podczas rozmieszczania modułów stosowana jest estymacja długości połączeń dla danego rozmieszczenia, ponieważ wyznaczanie ścieżek połączeń dla każdego przejściowego rozmieszczenia zajęłoby zbyt wiele czasu [4, 8]. Całkowita estymowana długość połączeń układu jest sumą estymowanych długości dla poszczególnych jego węzłów.

Po ustaleniu ostatecznego rozmieszczenia modułów z wykorzystaniem estymacji długości połączeń, wyznaczane są ścieżki połączeń. Czynność ta wykonywana jest jednak podczas następnego etapu projektowania topografii — podczas trasowania połączeń. Minimalne drzewo Steinera posiada tę własność, że łączy zadane punkty używając minimalnej długości połączeń. Do estymacji długości połączeń wykorzystywane są jednak aproksymacje minimalnego drzewa Steinera, ponieważ znalezienie tego drzewa w ogólnym przypadku jest problemem NP — trudnym. Najczęściej stosowanymi sposobami estymacji długości połączeń jest metoda połowy obwodu (ang. half-perimeter) oraz metoda wykorzystująca graf pełny (ang. complete graph) [1, 3–10, 27, 30, 38, 50–51]. Rysunek 9 przedstawia obydwie metody estymacji długości połączeń oraz drzewo Steinera — w przypadku grafu pełnego krawędzie grafu poprowadzono w przestrzeni Euklidesa, dla większej przejrzystości rysunku.



Rys. 9. Drzewo Steinera i metody estymacji długości połączeń: a) drzewo Steinera, b) połowa obwodu, c) graf pełny

Fig. 9. Steiner tree and wire length estimates: a) Steiner tree, b) half-perimeter, c) complete graph

Estymacja metodą połowy obwodu jest równa połowie obwodu prostokąta, obejmującego wszystkie końcówki danego węzła (kończówki modułów i końcówki całego układu). W metryce Manhattan ta estymacja jest równa minimalnej długości ścieżek niezbędnych do połączenia końcówek węzła, jeżeli liczba końcówek nie jest większa niż trzy. Dla większej liczby końcówek, wartość estymowana jest zwykle mniejsza od rzeczywistej minimalnej długości połączeń. Estymacja wykorzystująca graf pełny jest oparta na grafie pełnym, obejmującym wszystkie końcówki danego węzła, które stają się jego wierzchołkami. Graf pełny posiada krawędzie łączące wszystkie pary wierzchołków. Długość krawędzi grafu jest równa odległości między końcówkami w metryce Manhattan. Estymowana długość połączeń jest równa sumie długości wszystkich krawędzi grafu pełnego podzielonej przez $n/2$, gdzie n jest liczbą wszystkich końcówek w danym węźle. W grafie pełnym jest $n(n-1)/2$ wszystkich krawędzi, więc po podzieleniu przez $n/2$, otrzymujemy aproksymację długości $n-1$ połączeń, niezbędnych do połączenia n końcówek danego węzła. Wartość estymowana długości połączeń d jest określona wzorem

$$d = \frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n d_{ij} \quad (2)$$

gdzie: d_{ij} — długość połączenia między modułami i oraz j , we wzorze występuje dzielenie przez n , ponieważ długość każdej krawędzi grafu jest liczona dwukrotnie.

Dla obydwu metod możliwa jest korekcja estymowanej długości tak, aby odpowiadała rzeczywistej długości połączeń [1, 3–4, 9–10, 27].

4.3. MINIMALIZACJA OPÓŹNIEŃ W UKŁADZIE

Bardzo ważnym celem rozmieszczania modułów jest minimalizacja opóźnień w krytycznych częściach układu (ang. timing-driven placement). We współczesnych układach VLSI, opóźnienia wnoszone przez połączenia mają decydujące znaczenie na czas propagacji sygnałów, w porównaniu do opóźnień wnoszonych przez elementy aktywne. Jest to spowodowane stale postępującą miniaturyzacją układów oraz stałym wzrostem częstotliwości sygnału taktującego.

Optymalne rozmieszczenie modułów zapewnia znaczne zmniejszenie czasu propagacji sygnałów w układzie, co umożliwia zastosowanie w układzie sygnału taktującego o większej częstotliwości. Minimalizacja opóźnień w krytycznych częściach układu z reguły jest okupiona niewielkim wzrostem całkowitej długości połączeń. Metody rozmieszczania, które minimalizują opóźnienia w układzie, dzielimy na dwie grupy: metody oparte na ścieżkach (ang. path-based) oraz metody oparte na węzłach (ang. net based). Metody oparte na ścieżkach minimalizują opóźnienia w krytycznych ścieżkach, będących częścią różnych węzłów. Ze względu na bardzo dużą liczbę możliwych ścieżek, które mogą być rozpatrywane, metody te charakteryzują się bardzo dużą złożonością obliczeniową. Metody oparte na węzłach, minimalizują niezależnie opóźnienia w poszczególnych węzłach układu. Powszechnie stosowanym rozwiązaniem jest algorytm rozpatrujący „rezerwę” czasu na wyjściu wszystkich modułów (ang. zero-slack algorithm). W układzie musi być określony czas nadejścia sygnału na wejściach całego układu (ang. arrival time) oraz wymagany czas ustalenia się sygnału na jego wyjściach (ang. required time). W wyniku dwóch niezależnych przejść, od wejścia do wyjścia oraz od wyjścia do wejścia, określone są czasy nadejścia sygnału oraz wymagane czasy ustalenia sygnału dla wyjść wszystkich modułów. W przypadku czasu nadejścia sygnału brana jest pod uwagę najpóźniejsza jego zmiana. W przypadku wymaganego czasu ustalenia się sygnału brana jest pod uwagę najwcześniejsza konieczność jego zmiany. Uwzględniane są opóźnienia elementów aktywnych. Różnica między wymaganym czasem ustalenia się sygnału oraz czasem nadejścia sygnału określa „rezerwę” czasu dla danego modułu (ang. slack). Następnie algorytm przydziela opóźnienia do poszczególnych węzłów tak, aby „rezerwa” czasu dla wszystkich modułów była równa 0. Opóźnienia przydzielone poszczególnym węzłom są następnie zamieniane na wagi węzłów lub na ograniczenia ich maksymalnej długości. W celu określenia maksymalnej

długości węzłów wymagane są odpowiednie parametry technologiczne układu, które określają opóźnienia wnoszone przez połączenia w danym układzie. Zamiana opóźnień na wagi umożliwia rozwiązanie problemu minimalizacji opóźnień bez żadnych zmian w metodach rozmieszczania oraz bez dodatkowych nakładów obliczeniowych. Jest to możliwe, jeżeli funkcja kosztu danej metody wykorzystuje wagi poszczególnych węzłów układu. Węzły posiadające większą wagę będą wówczas preferowane podczas rozmieszczania. Możliwy jest taki dobór wag węzłów, który ułatwi dalszą minimalizację opóźnień, po wykonaniu rozmieszczenia modułów (ang. in-place optimization). Dalsza minimalizacja opóźnień w krytycznych częściach układu jest możliwa np. w wyniku zmiany rozmiarów poszczególnych tranzystorów, co zwiększa ich wydajność prądową [1, 3, 40–44].

4.4. METODY ROZMIESZCZANIA

Istnieje wiele metod rozmieszczania modułów w układach VLSI. Wyszukiwanie wyczerpujące (ang. exhaustive search) nie jest stosowane, ze względu na bardzo dużą liczbę możliwych rozwiązań. W przypadku układu składającego się z n części składowych o takich samych rozmiarach, liczba możliwych rozwiązań wynosi $n!$. Czasowa złożoność obliczeniowa wyszukiwania wyczerpującego jest zatem równa $O(n!)$. Wśród metod rozmieszczania można wyróżnić metody konstrukcyjne, iteracyjne, analityczne (numeryczne), wykorzystujące sieci neuronowe. Metody konstrukcyjne tworzą rozwiązanie poprzez umieszczanie w każdym kroku jednego modułu, w dogodnym dla niego miejscu. Metody konstrukcyjne umożliwiają bardzo szybkie otrzymanie rozwiązania, jednak jakość rozwiązań jest niezadowalająca [3–5, 9, 28]. Z tego względu obecnie metody te są stosowane do tworzenia początkowych rozwiązań dla metod iteracyjnych, ze względu na ich szybkość [3, 9]. Metody iteracyjne rozpoczynają swoje działanie od początkowego rozwiązania, które następnie poprawiają w kolejnych iteracjach. Wśród metod iteracyjnych należy wymienić: algorytm min-cut, metodę zamiany parami, symulowane wyżarzanie, algorytmy genetyczne, strategie ewolucyjne [1, 3–10, 20–23, 30]. Metody analityczne matematycznie wyznaczają położenie modułów [1, 3, 9, 30]. Inny sposób podziału metod rozmieszczania wyróżnia metody deterministyczne oraz stochastyczne. Metody deterministyczne rozwiązują dany problem rozmieszczenia w ten sam sposób podczas kolejnych prób [9].

5. PODSUMOWANIE

Niniejsza praca jest pierwszą częścią przeglądu, poświęconego metodom rozmieszczania w projektowaniu topografii cyfrowych układów VLSI. W pracy przedstawiono chronologicznie etapy procesu projektowania topografii. Opisano metody stosowane w poszczególnych etapach oraz wyjaśniono podstawowe pojęcia z tego zakresu. Szczegó-

łowo zostały przedstawione cele metod rozmieszczania oraz podano sposoby rozwiązania tego problemu.

Niniejsza praca stanowi zarys bardzo obszernej problematyki projektowania topografii układów VLSI. Metody rozmieszczania modułów zostaną szczegółowo opisane w kolejnych częściach przeglądu.

6. BIBLIOGRAFIA

1. M.J.S. Smith: *Application Specific Integrated Circuits*, Addison Wesley Longman, 1997.
2. A. Kos: *Modelowanie hybrydowych układów mocy i optymalizacja ich konstrukcji ze względu na rozkład temperatury*, Kraków, Wydawnictwa AGH, 1994.
3. B.T. Preas, M.J. Lorenzetti (red.): *Physical Design Automation of VLSI Systems*, Menlo Park, Benjamin Cummings, 1988.
4. T. Ohtsuki, (red.): *Layout Design and Verification*, Elsevier Science Publishers B. V. (North Holland), 1986.
5. M.M. Vai: *VLSI Design*, CRC Press, 2001.
6. C. Sechen: *VLSI Placement and Global Routing Using Simulated Annealing*, Boston, Kluwer Academic Publishers, 1988.
7. D.F. Wong, H.W. Leong, C. L. Liu: *Simulated Annealing for VLSI Design*, Kluwer Academic Publishers, 1988.
8. W. Wolf: *Modern VLSI Design: a systems approach*, Englewood Cliffs, New Jersey, PTR Prentice Hall, 1994.
9. K. Shahookar, P. Mazumder: *VLSI Cell Placement Techniques*, ACM Computing Surveys, 1991, vol. 23, pp. 143-220.
10. V. Betz, J. Rose: *VPR: A New Packing, Placement and Routing Tool for FPGA Research*, Proc. of the 7th International Workshop on Field Programmable Logic and Applications, London, 1997, pp. 213-222, <http://www.eecg.toronto.edu/~vaughn/papers/fpl97.pdf>.
11. J.G. Kim, Y.D. Kim: *A Linear Programming — Based Algorithm for Floorplanning in VLSI Design*, IEEE Transactions on Computer Aided Design, 2003, vol. 22, pp. 584-592.
12. M. Mogaki, C. Miura, H. Terai: *Algorithm for Block Placement with Size Optimization Technique by the Linear Programming Approach*, Proc. of the IEEE International Conference on Computer Aided Design, 1987, pp. 80-83.
13. Y.T. Lai, S.M. Leinwand: *Algorithms for Floorplan Design Via Rectangular Dualization*, IEEE Transactions on Computer Aided Design, 1988, vol. 7, pp. 1278-1289.
14. S. Wimer, I. Koren, I. Cederbaum: *Optimal Aspects Ratios of Building Blocks in VLSI*, IEEE Transactions on Computer Aided Design, 1989, vol. 8, pp. 139-145.
15. K. Chong, S. Sahni: *Optimal Realisations of Floorplans*, IEEE Transactions on Computer Aided Design, 1993, vol. 12, pp. 193-801.
16. G.K.H. Yeap, M. Sarrafzadeh: *A Unified Approach to Floorplan Sizing and Enumeration*, IEEE Transactions on Computer Aided Design, 1993, vol. 12, pp. 1858-1867.
17. W. Shi: *A Fast Algorithm for Area Minimization of Slicing Floorplans*, IEEE Transactions on Computer Aided Design, 1996, vol. 15, pp. 1525-1532.
18. P.S. Dasgupta, S. Sur-Kolay, B. Bhattacharya: *A Unified Approach to Topology Generation and Optimal Sizing of Floorplans*, IEEE Transactions on Computer Aided Design, 1998, vol. 17, pp. 126-135.
19. P. Pan, C. L. Liu: *Area Minimization for Floorplans*, IEEE Transactions on Computer Aided Design, 1995, vol. 14, pp. 123-132.

20. H. Murata, K. Fujiyoshi, S. Nakatake, Y. Kajitani: *VLSI Module Placement Based on Rectangle — Packing by the Sequence — Pair* IEEE Transactions on Computer Aided Design, 1996, vol. 15, pp. 1518-1524.
21. H. Murata, K. Fujiyoshi, M. Kaneko: *VLSI/PCB Placement with Obstacles Based on Sequence — Pair*, IEEE Transactions on Computer Aided Design, 1998, vol. 17, pp. 60-68.
22. K. Fujiyoshi, H. Murata: *Arbitrary Convex and Concave Rectilinear Block Packing Using Sequence — Pair*, IEEE Transactions on Computer Aided Design, 2000, vol. 19, pp. 224-233.
23. S. Nakatake, K. Fujiyoshi, H. Murata, Y. Kajitani: *Module Packing Based on the BSG — Structure and IC Layout Applications* IEEE Transactions on Computer Aided Design, 1998, vol. 17, pp. 519-530.
24. E.F.Y. Young, C.C.N. Chu, Z.C. Shen: *Twin Binary Sequences: A Nonredundant Representation for General Nonslicing Floorplan*, IEEE Transactions on Computer Aided Design, 2003, vol. 22, pp. 457-469.
25. Y. Feng, D.P. Mehta, H. Yang: *Constrained Floorplanning Using Network Flows*, IEEE Transactions on Computer Aided Design, 2004, vol. 23, pp. 572-580.
26. D. Sitaram, Y. Zheng, K.L. Shepard: *Full — Chip Three — Dimensional Shapes — Based RLC Extraction*, IEEE Transactions on Computer Aided Design, 2004, vol. 23, pp. 711-727.
27. A. Kos, Z. Nagórny: *Estymacja długości połączeń w układach VLSI*, IV Krajowa Konferencja Elektroniki, Dąbrowa Tarnobrzeg, czerwiec 2005, pp. 159-164.
28. M.A. Breuer (red.): *Automatyczne projektowanie maszyn cyfrowych*, Warszawa, PWN, 1976.
29. D.S. Rao, J.D. Provenç: *An integrated approach to routing and via minimization*, Information Processing Letters, 1991, vol. 39, pp. 257-263.
30. T.C. Hu, E.S. Kuh (red.): *VLSI Circuit Layout: Theory and Design*, New York, IEEE Press, 1985.
31. A.E. Dunlop, B.W. Kernighan: *A Procedure for Placement of Standard — Cell VLSI Circuits*, IEEE Transactions on Computer Aided Design, 1985, vol. 4, pp. 92-98.
32. P.R. Suaris, G. Kedem: *An Algorithm for Quadrissection and Its Application to Standard Cell Placement*, IEEE Transactions on Circuits and Systems, 1988, vol. 35, pp. 294-302.
33. A.E. Caldwell, A.B. Kahng, I.L. Markov: *Can Recursive Bisection Alone Produce Routable Placements?*, Design Automation Conference, 2000, pp. 477-482, <http://www.eecs.umich.edu/~imarkov/pubs/conf/c015.pdf>.
34. A.E. Caldwell, A.B. Kahng, I.L. Markov: *Optimal Partitioners and End Case Placers for Standard Cell Layout*, IEEE Transactions on Computer Aided Design, 2000, vol. 19, pp. 1304-1313.
35. M. Wang, X. Yang, M. Sarrafzadeh: *Dragon2000: Standard — Cell Placement Tool for Large Industry Circuits*, International Conference on Computer Aided Design, 2000, pp. 260-263, <http://er.cs.ucla.edu/Dragon/papers/dragon.pdf>.
36. A.R. Agnihotri, S. Ono, C. Li, M.C. Yildiz, A. Khatkhate, C.K. Koh, P.H. Madden: *Mixed Block Placement via Fractional Cut Recursive Bisection*, IEEE Transactions on Computer Aided Design, 2005, vol. 24, pp. 748-761.
37. M.C. Yildiz, P.H. Madden: *Improved Cut Sequences for Partitioning Based Placement*, Design Automation Conference, 2001, pp. 776-779, <http://vlsicad.cs.binghamton.edu/pubs/Yildiz01dac.pdf>.
38. P.H. Madden: *Reporting of Standard Cell Placement Results*, IEEE Transactions on Computer Aided Design, 2002, vol. 21, pp. 240-247.
39. S.N. Adya, M.C. Yildiz, I.L. Markov, P.G. Villarrubia, P.N. Parakh, P.H. Madden: *Benchmarking for Large — Scale Placement and Beyond*, IEEE Transactions on Computer Aided Design, 2004, vol. 23, pp. 472-486.
40. X. Yang, B.K. Choi, M. Sarrafzadeh: *Timing — Driven Placement using Design Hierarchy Guided Constraint Generation*, International Conference on Computer Aided Design, 2002, pp. 177-184, <http://er.cs.ucla.edu/Dragon/papers/timingdragon.pdf>.
41. R. Nair, C.L. Berman, P.S. Hauge, E.J. Yoffa: *Generation of Performance Constraints for Layout*, IEEE Transactions on Computer Aided Design, 1989, vol. 8, pp. 860-874.

42. H. Youssef, R.B. Lin, E. Shragowitz: *Bounds on Net Delays for VLSI Circuits*, IEEE Transactions on Circuits and Systems, 1992, vol. 39, pp. 815-824.
43. H. Ren, D.Z. Pan, D.S. Kung: *Sensitivity Guided Net Weighting for Placement — Driven Synthesis*, IEEE Transactions on Computer Aided Design, 2005, vol. 24, pp. 711-721.
44. Q. Liu, B. Hu, M. Marek Sadowska: *Individual Wire Length Prediction With Application to Timing — Driven Placement*, IEEE Transactions on VLSI Systems, 2004, vol. 12, pp. 1004-1014.
45. T. Łuba, K. Jasiński, B. Zbierzchowski: *Specjalizowane układy cyfrowe w strukturach PLD i FPGA*, Warszawa, Wydawnictwa Komunikacji i Łączności, 1997.
46. T. Łuba, B. Zbierzchowski: *Komputerowe projektowanie układów cyfrowych*, Warszawa, Wydawnictwa Komunikacji i Łączności, 2000.
47. S. Sutanthavibul, E. Shragowitz, J.B. Rosen: *An Analytical Approach to Floorplan Design and Optimization*, IEEE Transactions on Computer Aided Design, 1991, vol. 10, pp. 761-769.
48. H.H. Chan, I.L. Markov: *Practical Slicing and Non — slicing Block — Packing without Simulated Annealing*, Proc. of the Great Lakes Symposium on VLSI, 2004, pp. 282-287, <http://vlsicad.eecs.umich.edu/BK/BloBB/PAPERS/p037-chan.pdf>.
49. F. Mo, A. Tabbara, R.K. Brayton: *A Force — Directed Macro — Cell Placer*, Proc. of the International Conference on Computer Aided Design, 2000, pp. 177-180, <http://www.sigda.org/Archives/ProceedingArchives/Iccad/>.
50. R. Sedgewick: *Algorithms in C, Part 5: Graph Algorithms*, Addison Wesley Professional, 2002.
51. R. Sedgewick: *Algorytmy w C++. Część 5. Grafy*, Warszawa, Wydawnictwo RM, 2003.

Z. NAGÓRNY, A. KOS

THE DESIGN OF THE VLSI CIRCUIT LAYOUT — PART 1. STYLES, PHASES, PLACEMENT

Summary

The design process of the VLSI circuits requires the use of computer aided design tools. This paper is the first part of the survey of the cell placement techniques for digital VLSI circuits. Design styles used in VLSI circuits are described. Layouts of Standard Cell, Gate Array, Sea-of-Gates and Field Programmable Gate Array are presented. Then, the physical design flow, which includes partitioning, floorplanning, placement, routing and verification are described. Special attention is paid to floorplanning, because this stage is similar to the placement problem. The cell placement problem and placement techniques are described. VLSI cell placement is a very important phase of the physical design process. Cell placement, which is a very difficult optimization problem, has proved to be a NP — compete. The goal of the VLSI cell placement is to arrange all the cells on a placement carrier while minimizing an objective or cost function. The most commonly used objectives of the placement are to minimize the total estimated wire length and the interconnect congestion, and to meet the timing requirements for critical nets. Commonly used wire length estimates for the cell placement are presented. The timing driven placement methods are described. The algorithms used for the cell placement are presented.

Keywords: VLSI, physical design process, integrated circuit, layout, computer aided design, Standard Cell, Gate Array, Sea-of-Gates, Field Programmable Gate Array, partitioning, floorplanning, cell placement, routing, wire length estimation, half-perimeter, complete graph, timing-driven placement, constructive placement methods, min-cut algorithm, pairwise interchange algorithm, simulated annealing, genetic algorithm, analytical methods, neural network

Projektowanie topografii systemów VLSI Cz. 2. Algorytm min-cut

ZBIGNIEW NAGÓRNY, ANDRZEJ KOS

*Zbigniew Nagórny,
Studium Generale Sandomiriense,
Wyższa Szkoła Humanistyczno-Przyrodnicza w Sandomierzu,
ul. Krakowska 26, 27-600 Sandomierz,
e-mail: nagorny@interia.pl,*

*Andrzej Kos,
Akademia Górniczo-Hutnicza w Krakowie, Katedra Elektroniki,
al. Mickiewicza 30, 30-059 Kraków,
e-mail: kos@agh.edu.pl*

*Otrzymano 2005.10.06
Autoryzowano 2006.02.20*

Niniejsza praca jest drugą częścią przeglądu metod rozmieszczania modułów, stosowanych podczas projektowania topografii układów VLSI. W pracy szczegółowo został opisany algorytm min-cut. Przedstawiono algorytm Kernighana i Lina, który jest stosowany w algorytmie min-cut. Opisano algorytm podziału Fiduccia i Mattheysesa. Przedstawiono modyfikacje algorytmu min-cut. Podany został sposób zastosowania algorytmu min-cut dla topografii swobodnej. Omówiono wielopoziomowy algorytm podziału hMETIS. Scharakteryzowano obecnie stosowane programy, które wykorzystują algorytm min-cut: Capo, Dragon, Feng Shui, QUAD.

Słowa kluczowe: VLSI, projektowanie topografii układu, układ scalony, projektowanie wspomaganie komputerowo, układ Standard Cell, rozmieszczanie modułów, podział układu, algorytm Kernighana i Lina, algorytm Fiduccia i Mattheysesa, algorytm min-cut, wielopoziomowy podział układu, algorytm hMETIS, Capo, Dragon, Feng Shui, QUAD

1. WPROWADZENIE

Niniejsza praca jest drugą częścią przeglądu metod rozmieszczania modułów, stosowanych podczas projektowania topografii układów VLSI. W pracy został szczegóło-

wo opisany algorytm min-cut oraz programy wykorzystujące ten algorytm. Algorytm min-cut został zaproponowany przez Breuera w 1977r. [10]. W ciągu minionych lat był wielokrotnie modyfikowany. Obecnie nadal jest powszechnie stosowany, a w ostatnich latach pojawiło się wiele programów rozmieszczania, które wykorzystują ten algorytm [11, 20, 24–29, 31–34]. Algorytm min-cut wykorzystuje metodę „dziel i rządź”, która umożliwia podział problemu na mniejsze podproblemy. Do podziału układu stosowany jest znany heurystyczny algorytm Kernighana i Lina lub jego modyfikacje. W rozdziale 2 przedstawiono algorytm Kernighana i Lina oraz jego modyfikacje. Algorytm min-cut jest opisany w rozdziale 3. Wielopoziomowy algorytm podziału przedstawiono w rozdziale 4. Rozdział 5 zawiera charakterystykę obecnie stosowanych programów, których podstawą jest algorytm min-cut. Podsumowanie drugiej części przeglądu metod rozmieszczania zamieszczono w rozdziale 6. Bibliografia znajduje się w rozdziale 7.

2. ALGORYTM KERNIGHANA I LINA

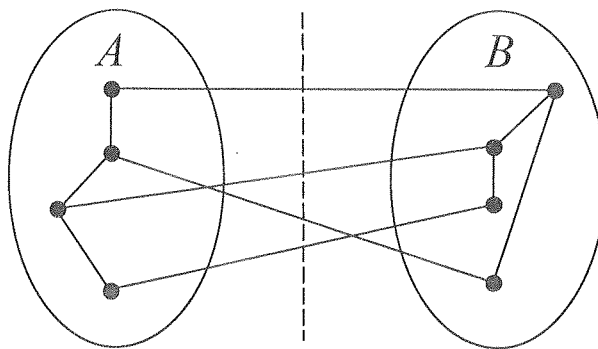
Algorytm Kernighana i Lina (ang. group migration method) stosowany jest do podziału układu elektronicznego na części tak, aby koszt połączeń między tymi częściami był minimalny. Algorytm Kernighana i Lina jest stosowany w projektowaniu topografii układu VLSI podczas etapu podziału [1, 3–4]. Stosowany jest również w algorytmie min-cut [1, 3, 6, 9]. Problem podziału układu można przedstawić następująco. Dany jest graf S , który odzwierciedla połączenia między $2n$ modułami w układzie VLSI. Wierzchołkami grafu są moduły, natomiast krawędziami grafu są połączenia między modułami. Połączenia w układzie są określone przez macierz połączeń $C = [c_{ij}]$, dla $i, j = 1, \dots, 2n$. Element macierzy c_{ij} określa koszt (wagę) połączenia między modułami i oraz j . Macierz C jest macierzą symetryczną oraz $c_{ii} = 0$, dla $i = 1, \dots, 2n$. Ponadto zakłada się, że wszystkie węzły układu zawierają tylko dwie końcówki. Należy znaleźć podział grafu S na dwie równe części A i B tak, aby suma kosztów połączeń między modułami znajdującymi się w różnych podzbiorach była minimalna. Rysunek 1 ilustruje problem podziału układu.

Założmy, że graf S został arbitralnie podzielony na dwie równe części A i B , po n modułów tak, że $S = A \cup B$. Dla każdego $a \in A$ określamy koszt zewnętrzny E_a oraz koszt wewnętrzny I_a , w następujący sposób

$$E_a = \sum_{y \in B} c_{ay} \quad (1)$$

$$I_a = \sum_{x \in A} c_{ax} \quad (2)$$

Podobnie określamy koszt zewnętrzny oraz wewnętrzny dla każdego $b \in B$. Koszt wewnętrzny określa sumę kosztów połączeń danego modułu z modułami położonymi w tym samym podzbiorze, koszt zewnętrzny dotyczy połączeń z modułami znajdującymi



Rys. 1. Problem podziału układu

Fig. 1. Circuit partitioning problem

się w drugim podzbiorze. Jeżeli zamienimy miejscami $a \in A$ oraz $b \in B$, to można pokazać, że zmiana kosztu g (ang. gain) określona jest wzorem [1, 10]

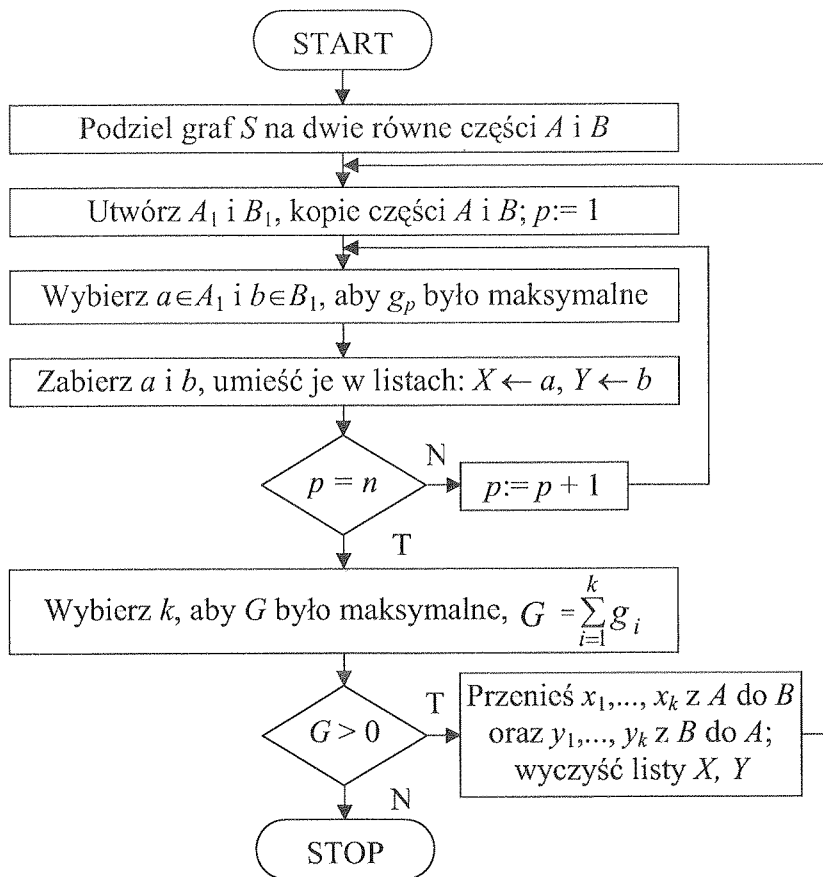
$$g = \text{stary koszt} - \text{nowy koszt} = (E_a - I_a) + (E_b - I_b) - 2c_{ab} \quad (3)$$

W celu rozwiązania problemu podziału, graf S należy najpierw podzielić na dwie równe części A i B , w dowolny sposób. Tworzone są kopie obydwu części, odpowiednio A_1 i B_1 . Następnie z obydwu części wybieramy $a \in A_1$ oraz $b \in B_1$, dla których zmiana kosztu jest maksymalna. Moduły a i b zabieramy z obydwu części i umieszczamy je w dwóch listach X i Y , odpowiednio dla części A_1 oraz B_1 . Powyższe czynności powtarzamy n razy, oznaczając maksymalne zmiany kosztu dla poszczególnych par przez g_p , dla $p = 1, \dots, n$. W ten sposób w listach zostaną umieszczone wszystkie moduły. Wprowadza się sumaryczną zmianę kosztu G , która jest równa sumie zmian kosztów g_p , odpowiadających pierwszym k parom modułów w listach

$$G = \sum_{i=1}^k g_i \quad (4)$$

Następnie znajdujemy taką wartość k , aby sumaryczna zmiana kosztu G była maksymalna, k może przyjmować wartości z przedziału $1, \dots, n$. Następnie, w oryginalnych częściach A i B zamieniamy miejscami k pierwszych par modułów z utworzonych list. Powyższe czynności powtarzamy dopóki G jest większe od 0, ponieważ następuje wtedy zmniejszenie kosztu połączeń. Rysunek 2 przedstawia algorytm postępowania [1, 3–4, 9–10, 29].

Algorytm Kernighana i Lina nie umożliwia podziału obwodu, którego węzły posiadają więcej niż dwie końcówki. W celu podziału układu na nierówne części, konieczne jest stosowanie tzw. „sztucznych” modułów (ang. dummy cell), które nie posiadają żadnych połączeń. W dodatku wszystkie moduły muszą posiadać taki sam rozmiar i nie można trwale narzucić ich położenia przed podziałem. Algorytm Kernighana i Lina posiada czasową złożoność obliczeniową $O(n^2 \ln n)$ [1, 6, 10, 13, 29].



Rys. 2. Algorytm Kernighana i Lina

Fig. 2. Flow chart of the Kernighan-Lin algorithm

Schweikert i Kernighan zmodyfikowali algorytm podziału, aby umożliwić jego stosowanie w przypadku, gdy węzły układu posiadają więcej niż dwie końcówki [1, 3, 6, 9–10]. Koszt podziału układu jest równy liczbie węzłów układu, które przecina linia jego podziału. Zastosowany został inny sposób obliczenia zmiany kosztu dla zamiany miejscami pary modułów [8, 10, 13]

$$g = \text{stary koszt} - \text{nowy koszt} = D_a + D_b - \text{korekta}_{ab} \quad (5)$$

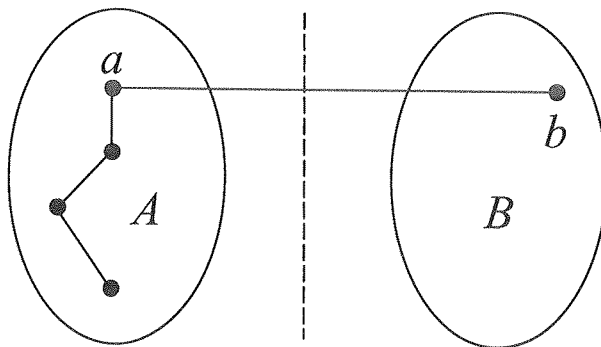
gdzie: D_a – zmiana kosztu po przeniesieniu $a \in A$ do części B , D_b – zmiana kosztu po przeniesieniu $b \in B$ do części A , korekta_{ab} – korekta zmiany kosztu dla węzła zawierającego obydwa moduły a i b .

Zmiana kosztu podziału jest sumą zmian kosztu dla poszczególnych węzłów układu. Wykorzystana jest właściwość łączenia modułów w węzle, który został podzielony na dwie części. Do połączenia modułów znajdujących się w dwóch częściach układu

wystarczy tylko jedno połączenie między tymi częściami. Koszt podziału ze względu na ten węzeł jest wtedy równy 1. W związku z tym, zmiana kosztu D_a , ze względu na dany węzeł wynosi:

- 1, jeżeli w części A układu, moduł a jest jedynym modulem należącym do danego węzła
- 0, jeżeli w części A jest więcej niż jeden moduł, natomiast w części B jest przynajmniej jeden moduł należący do danego węzła
- 1, jeżeli wszystkie moduły danego węzła znajdują się w części A

Wielkość korekty dla danego węzła $korekta_{ab}$ wynosi 1, jeżeli obydwa moduły a i b należą do tego węzła i jeden z modułów a lub b jest jedynym modulem danego węzła, w jednej z części układu. Rysunek 3 przedstawia powyższą sytuację. W rozpatrywanej sytuacji, zmiana kosztu dla modułu a jest równa 0, a dla modułu b wynosi 1. Zamiana miejscami modułów a i b nie powoduje jednak w tym przypadku zmiany kosztu połączeń. W wyniku zastosowania korekty, zmiana kosztu dla tego przypadku jest prawidłowo obliczana i wynosi 0. Dla pozostałych przypadków $korekta_{ab}$ wynosi 0 [8, 10].



Rys. 3. Korekta zmiany kosztu

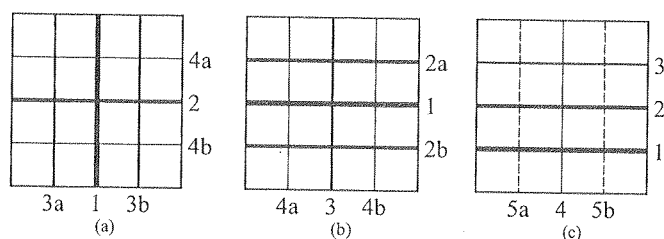
Fig. 3. Correction of the gain

Fiduccia i Mattheyses, a później Krishnamurthy, zmodyfikowali algorytm podziału Kernighana i Lina. Moduły mogą posiadać dowolne rozmiary oraz trwale ustalone położenie w układzie. Za każdym razem tylko jeden moduł zmienia swoje położenie. Algorytm kontroluje sumaryczną powierzchnię modułów w obydwu częściach powstających w wyniku podziału, którym przydzielana jest cała powierzchnia dzielonego układu. Części układu mogą posiadać różną powierzchnię. Suma powierzchni modułów w każdej części nie może przekroczyć przydzielonej powierzchni (ang. balance). Zmiana kosztu, odpowiadająca danemu modułowi, decyduje o jego przeniesieniu do drugiej części. Przenoszone są moduły, którym odpowiada maksymalna zmiana kosztu, pod warunkiem, że nie narusza to ustalonego przydziału powierzchni dla poszczególnych części. Po przeniesieniu do drugiej części, moduł jest blokowany i nie jest dalej rozpatrywany. Po przeniesieniu wszystkich możliwych modułów, odtwarzany jest

podział, któremu odpowiada najmniejszy koszt, w całej serii przenoszenia modułów (ang. pass). Serie przenoszenia modułów są powtarzane do momentu, gdy nie następuje już dalsze zmniejszanie kosztu. Po każdym przeniesieniu modułu aktualizowane są jedynie zmiany kosztu dla modułów, które należą do tych węzłów układu, do których należał ostatnio przeniesiony moduł. Taki sposób aktualizacji zmiany kosztu zapewnia większą efektywność algorytmu w porównaniu do algorytmu Kernighana i Lina. W przypadku, gdy koszt przeniesienia jest identyczny dla kilku modułów, mogą być rozpatrywane koszty przeniesienia modułów w następnym kroku. Preferowane są moduły, które umożliwią większą zmianę kosztu w kolejnych krokach. Technika ta jest określana terminem lookahead (ang.). Algorytm Fiduccia i Mattheysesa posiada złożoność obliczeniową $O(n)$ [1, 3, 6, 12, 15–18, 21–22].

3. ALGORYTM MIN-CUT

Algorytm min-cut opiera się na podziale układu na części tak, aby koszt połączeń między tymi częściami był minimalny. Części powstające w wyniku podziału są dalej dzielone, dopóki każda z nich nie będzie zawierać dokładnie jednego modułu. W ten sposób każdy moduł ma określone położenie w układzie. Podczas podziału moduły są przenoszone między częściami układu, z wykorzystaniem algorytmu Kernighana i Lina lub jego modyfikacji [1, 3–6, 8–10]. Breuer zaproponował trzy sposoby podziału układu: podział na cztery części (ang. quadrature placement), podział na dwie połowy (ang. bisection placement) oraz podział typu slice/bisection (ang.). Podział na cztery części polega na ciągłym podziale układu na przemian liniami pionowymi i poziomymi. Ten sposób podziału jest najkorzystniejszy dla układów, które posiadają duże skupienie połączeń w środkowej części układu. Pierwszy podział układu na cztery części umożliwia zmniejszenie skupienia połączeń w środku układu. Podział na cztery części jest najczęściej stosowanym sposobem podziału. Podział na dwie połowy polega na ciągłym podziale układu na dwie równe połowy, najpierw liniami poziomymi, a następnie pionowymi. Podział liniami poziomymi umożliwia podzielenie układu na wiersze, do których przypisane są określone moduły. Następnie, podział liniami pionowymi umożliwia wyznaczenie pozycji modułów w wierszach. Ten rodzaj podziału jest korzystny dla układów Standard Cell. Podział typu slice/bisection polega na podziale układu najpierw liniami poziomymi, które mogą dzielić układ na części o różnych rozmiarach. Podział liniami poziomymi umożliwia przypisanie modułów do wierszy. Następnie układ jest dzielony liniami pionowymi na dwie równe części, co umożliwia określenie położenia modułów w wierszach. Ten sposób podziału jest korzystny dla układów, które posiadają duże skupienie połączeń na obrzeżach układu [3–4, 9–10]. Należy podkreślić, że po ustaleniu położenia modułów względem danej linii podziału, położenie modułów względem tej linii nie może ulec zmianie (moduły nie mogą być przeniesione na drugą stronę linii). Rysunek 4 przedstawia stosowane sposoby podziału układu, kolejność podziału określa numer oraz szerokość linii. Dla każdego sposobu, kolejność podziału liniami poziomymi i pionowymi może być odwrócona.

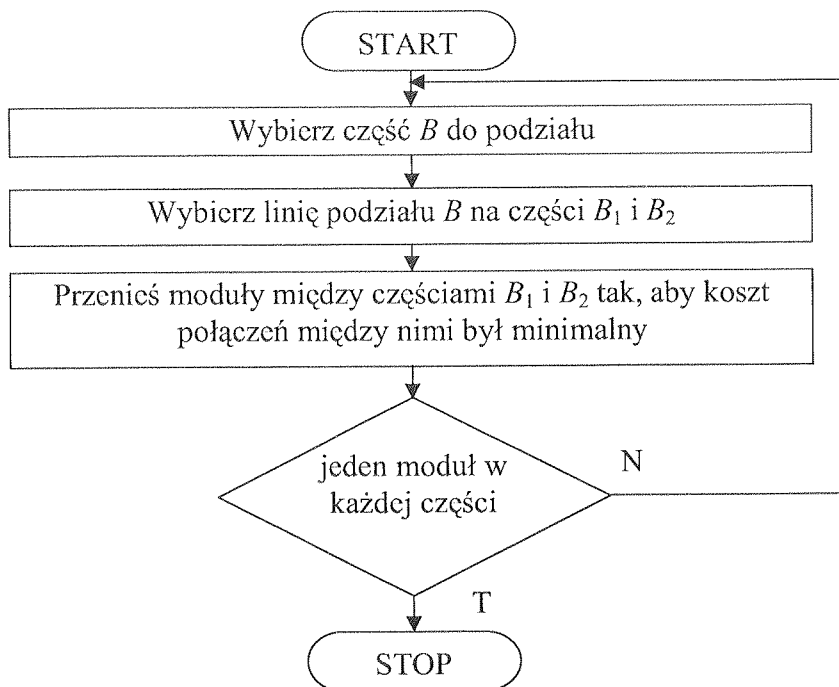


Rys. 4. Sposoby podziału układu w algorytmie min-cut: a) podział na cztery części, b) podział na dwie połowy, c) podział typu slice/bisection

Fig. 4. Cut sequences used in the min-cut algorithm: a) quadrature placement, b) bisection placement, c) slice/bisection placement

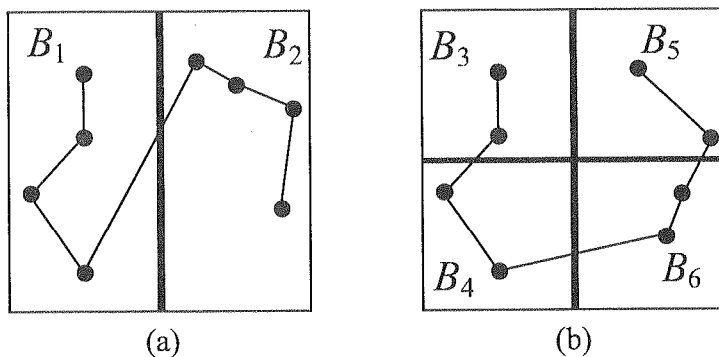
Rysunek 5 przedstawia sposób postępowania w algorytmie min-cut. Każdej części układu odpowiada określony poziom, poczynając od poziomu 1 dla całego układu. W wyniku podziału danej części układu, której odpowiada poziom j , powstają części z przypisanym poziomem $j + 1$. Poszczególne części układu mogą być dzielone w różnej kolejności. Z reguły najpierw są dzielone wszystkie bloki posiadające przypisany poziom j , a następnie są dzielone części na poziomie $j + 1$. Sposób ten jest podziałem wszerz (ang. breadth first). W drugim sposobie, w podziale w głąb (ang. depth first), do następnego podziału wybierana jest część, która może być dalej dzielona i posiada największy przypisany poziom. Części układu, które powstają w wyniku jego podziału, są kolejno oddzielnie przetwarzane. Umożliwia to znaczną redukcję nakładów obliczeniowych, jednak nie zapewnia osiągnięcia minimum globalnego funkcji kosztu. Korzystne może być również stosowanie odrębnych linii podziału dla poszczególnych części układu. Początkowe położenie modułów w układzie może być określone w sposób losowy. W tym celu można zastosować również algorytm konstrukcyjny, który może wykorzystywać ustaloną pozycję niektórych modułów w układzie [9–10].

Rysunek 6 przedstawia podział układu na cztery części: B_3 , B_4 , B_5 , B_6 . W powyższym przykładzie, wszystkie połączenia stanowią odrębne węzły. Breuer zaproponował iteracyjny sposób podziału układu. Po wykonaniu podziału układu linią pionową na części B_1 i B_2 , każda z tych części jest dalej dzielona linią poziomą. Podczas podziału części B_1 ignorowane są moduły znajdujące się w części B_2 , ponieważ nie jest określone ich położenie względem linii poziomej. Podczas podziału B_2 na części B_5 i B_6 , uwzględniane jest położenie modułów znajdujących się w częściach B_3 i B_4 , ponieważ ma ono wpływ na koszt podziału B_2 linią poziomą. Po podziale części B_2 , korygowany jest podział części B_1 . W tym celu wykorzystywane jest ustalone już położenie modułów w częściach B_5 i B_6 , ponieważ ma ono wpływ na koszt podziału B_1 linią poziomą. Następnie, na tej samej zasadzie odbywa się korekcja podziału części B_2 . Iteracyjna korekcja podziału części B_1 oraz B_2 jest wykonywana na przemian, aż do ustalenia wyniku. Podczas podziału układu w algorytmie zaproponowanym przez Breuera, nie jest jednak uwzględniane położenie modułów w pozostałych częściach układu.



Rys. 5. Algorytm min-cut

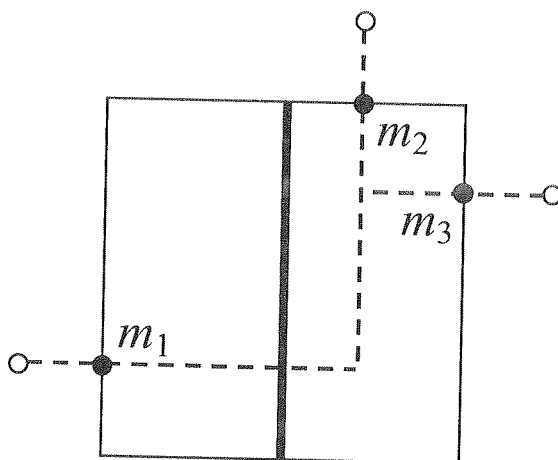
Fig. 5. Flow chart of the min-cut algorithm



Rys. 6. Iteracyjny podział układu w algorytmie min-cut: a) podział linią pionową, b) podział linią poziomą

Fig. 6. Iterative circuit partitioning in the min-cut algorithm: a) partitioning by a vertical cut line, b) partitioning by a horizontal cut line

Powyższy problem rozwiązuje modyfikacja wprowadzona przez Dunlopa i Kernighana (ang. terminal propagation). Podczas podziału układu lub jego części, są uwzględniane połączenia z modułami znajdującymi się w innych częściach oraz połączenia z końcówkami całego układu. Zakłada się, że moduły umieszczone w innych częściach układu są położone w środku geometrycznym tych części. W ogólnym przypadku, dla każdego węzła układu, moduły położone a dzielonej części, końcówki całego układu oraz moduły położone na zewnątrz dzielonej jego części, są łączone z użyciem minimalnego drzewa Steinera. Punkty przecięcia tego drzewa z granicami dzielonej części, są traktowane jak dodatkowe moduły, o trwale ustalonej pozycji. Dodatkowe moduły położone blisko linii podziału nie są uwzględniane. W ten sposób, moduły znajdujące się w pozostałych częściach układu mają wpływ na sposób podziału danej jego części [13]. Rysunek 7 ilustruje opisaną modyfikację. Na obrzeżach dzielonej części układu zostały umieszczone dodatkowe moduły m_1 , m_2 i m_3 .



Rys. 7. Modyfikacja Dunlopa i Kernighana

Fig. 7. Terminal propagation

Powyższa modyfikacja umożliwia zmniejszenie rozmiarów układu o ok. 30%, w porównaniu do wyników uzyskanych z użyciem zwykłego algorytmu min-cut. Technika ta wymaga podziału układu wszęsz [6, 9, 13]. Położenie dodatkowych modułów może być również określone bez konieczności konstruowania minimalnego drzewa Steinera. Dodatkowe moduły mogą być umieszczane w najbliższych punktach obrzeża dzielonej części układu, mogą nie być uwzględniane w przypadku, gdy znajdują się blisko linii podziału [6, 9, 13].

Suaris i Kedem wprowadzili podział każdej części układu równocześnie na cztery ćwiartki (ang. quadrissection). Do podziału układu został zastosowany algorytm Fiduccia i Mattheyses. Funkcja kosztu jest określona przez topologię połączeń między modułami węzła, znajdującymi się w różnych ćwiartkach dzielonej części układu. Al-

gorytm pozwala na stosowanie różnych wag połączeń dla kierunku poziomego i pionowego, co może być wykorzystane do zmniejszenia wysokości kanałów w układach Standard Cell. Uwzględniane są połączenia z modułami znajdującymi się poza dzieloną częścią układu. W tym celu stosowana jest modyfikacja Dunlopa i Kernighana. Podczas podziału danej części stosowana jest iteracyjna korekcja jej podziału. Korekta jest wykonywana cyklicznie dla wszystkich ćwiartek dzielonej części, aż do momentu ustalenia wyniku. Wymaga to zwykle kilku korekt podziału dla każdej z ćwiartek [9, 14]. Algorytm ten stał się podstawą działania programu QUAD [11, 25, 29].

Algorytm min-cut może być stosowany w układach o topografii swobodnej, w której moduły posiadają dowolne rozmiary i położenie [3–4, 6, 9]. Do określenia wzajemnego położenia modułów stosowane są dwa planarne, acykliczne grafy skierowane, odpowiednio dla kierunku poziomego i pionowego. Układ jest dzielony na części liniami poziomymi i pionowymi tak, aby koszt połączeń między nimi był minimalny oraz różnica między sumarycznymi powierzchniami modułów w powstających częściach, nie przekraczała ustalonej wartości. Rozmiary modułów nie są brane pod uwagę. Podczas podziału danej części, stosowana jest iteracyjna korekcja jej podziału. Podział układu jest kontynuowany do momentu, kiedy w każdej części będzie znajdował się jeden moduł. W ten sposób ustalone jest wzajemne położenie modułów, które jest podstawą do ostatecznego określenia ich położenia, po uwzględnieniu rozmiarów modułów [4].

Algorytm min-cut oparty jest na minimalizacji skupienia połączeń, która z kolei pośrednio wpływa na minimalizację powierzchni układu, ze względu na mniejsze rozmiary kanałów do prowadzenia połączeń. Minimalizacja skupienia połączeń pośrednio przyczynia się do minimalizacji długości połączeń, ponieważ silnie połączone ze sobą moduły są rozmieszczane obok siebie [1, 5, 8–9, 13, 20, 23–24, 26]. W rozdziale 5 zostaną przedstawione programy rozmieszczania, które wykorzystują algorytm min-cut oraz jego modyfikacje.

4. WIELOPOZIOMOWY PODZIAŁ UKŁADU

Algorytm Fiduccia i Mattheysesa był wielokrotnie modyfikowany [16–17]. Przełomowym osiągnięciem dla algorytmu podziału oraz algorytmu min-cut jest jednak wprowadzenie wielopoziomowego podziału układu (ang. multilevel circuit partitioning) [16–18]. Metoda ta umożliwia podział układów o bardzo dużych rozmiarach, rzędu setek tysięcy modułów.

Silnie połączone moduły są najpierw ze sobą łączone. Następnie tworzona jest nowa lista połączeń układu, w której połączone moduły są reprezentowane przez jeden moduł (klastery) o odpowiednio większej powierzchni (ang. cluster). Powyższe czynności są powtarzane wiele razy, w rezultacie następuje zmniejszanie liczby modułów w liście połączeń (ang. coarsening phase). Każdej nowej liście połączeń odpowiada odpowiedni, wyższy poziom. Listy połączeń na wszystkich poziomach są zapisywane. Powyższe czynności powtarzane są do momentu, kiedy lista połączeń zawiera odpowiednio małą liczbę modułów. Następnie odbywa się wstępny podział układu (grafu

odzwierciedlającego połączenia w układzie) na dwie części, który może być wykonany w sposób losowy. Następnie, wykonywane są podziały na kolejnych, niższych poziomach, ale zgodnie z listą połączeń rozpatrywanego poziomu (ang. uncoarsening and refinement phase). Podział układu kończy się po podziale układu opisywanego przez oryginalną listę połączeń. Rezultat podziału na wyższym poziomie jest zawsze wykorzystywany jako początkowy podział na kolejnym, niższym poziomie. Na każdym poziomie rozdzielane są moduły, które wcześniej zostały połączone na tym poziomie, podczas etapu łączenia modułów. Celem podziału na każdym poziomie jest osiągnięcie minimalnego kosztu połączeń między dwoma częściami układu. Podział układu na kolejnych poziomach może być wykonywany z wykorzystaniem zwykłego algorytmu podziału, np. algorytmu Fiduccia i Mattheyses. Wynik podziału układu na wyższym poziomie jest bardzo dobrym początkowym podziałem, dla podziału na kolejnym, niższym poziomie. Dzięki temu algorytm wykonuje znacznie mniejszą liczbę przeniesień modułów w porównaniu do tradycyjnego rozwiązania. Mniejsza liczba przeniesień modułów ma decydujący wpływ na dużą efektywność algorytmu, nawet dla bardzo dużych układów [16–18].

4.1. ALGORYTM HMETIS

Powszechnie stosowanym algorytmem wielopoziomowego podziału jest algorytm hMETIS. W algorytmie stosowane są trzy sposoby łączenia modułów. Pierwszy sposób polega na łączeniu par modułów, które należą do węzła zawierającego tylko dwie końcówki (ang. edge coarsening). Moduły są rozpatrywane kolejno w porządku losowym. Dla każdego rozpatrywanego modułu sprawdzane są wszystkie moduły, które wraz z rozpatrywanym modułem tworzą węzeł, składający się tylko z dwóch końcówek. Waga odpowiadająca takiej parze połączonych modułów jest równa sumie wag odpowiadających wszystkim węzłom układu, które łączą daną parę modułów w układzie. Waga odpowiadająca węzłowi łączącemu daną parę modułów jest równa $1/(n - 1)$, gdzie n oznacza liczbę końcówek w węźle. Kojarzone są pary, którym odpowiada największa waga, pod warunkiem, że żaden z modułów nie został wcześniej skojarzony. Po rozpatrzeniu wszystkich modułów, łączone są wszystkie skojarzone pary modułów. Drugi sposób polega na łączeniu modułów, które należą do węzła posiadającego dowolną liczbę końcówek (ang. hyperedge coarsening). Węzły są sortowane według nierosnącej wagi. Węzły, którym odpowiada taka sama waga, sortowane są według niemalejącej liczby końcówek w węźle. Następnie, węzły są rozpatrywane zgodnie z ustalonym porządkiem. Kojarzone ze sobą są wszystkie moduły w danym węźle, pod warunkiem, że żaden z modułów nie został wcześniej skojarzony. Po rozpatrzeniu wszystkich węzłów, łączone są wszystkie skojarzone grupy modułów. Drugi sposób zapewnia połączenie znacznie większej liczby modułów, w porównaniu do pierwszego rozwiązania. Duża część modułów nie jest jednak łączona. Ma to miejsce w przypadku, gdy część modułów należących do danego węzła, została wcześniej skojarzona podczas rozpatrywania innego węzła. Trzeci sposób rozwiązuje ten problem (ang. modified hyperedge

coarsening). Po połączeniu wszystkich modułów zgodnie z drugim sposobem, lista węzłów jest rozpatrywana ponownie. Łączone są ze sobą wyłącznie te moduły spośród modułów należących do poszczególnych węzłów, które nie zostały wcześniej ze sobą połączone [17].

Początkowy podział układu na najwyższym poziomie odbywa się w sposób losowy. Wskazane jest utworzenie kilku początkowych podziałów układu, które są następnie niezależnie rozpatrywane. Zapewnia to uzyskanie lepszych wyników. Podczas podziału układu na kolejnych poziomach, stosowany jest algorytm Fiduccia i Mattheysesa oraz algorytm podziału, który przenosi równocześnie kilka modułów (ang. hyperedge refinement). W algorytmie Fiduccia i Mattheysesa zastosowano dwie modyfikacje. Pierwsza polega na wykonywaniu tylko dwóch serii przenoszenia modułów podczas każdego podziału. Druga modyfikacja polega na przerywaniu danej serii przenoszenia modułów w przypadku, gdy przeniesienie k kolejnych modułów ($k = 1\%$ wszystkich modułów w liście połączeń), nie powoduje poprawy kosztu podziału. Wprowadzone modyfikacje umożliwiają znaczne przyspieszenie działania algorytmu, bez istotnego pogorszenia jakości otrzymywanych wyników. Drugi z algorytmów podziału może przenosić równocześnie kilka modułów należących do określonego węzła. Przenoszenie modułów ma na celu umieszczenie wszystkich modułów należących do danego węzła w jednej części układu [17].

Istnieje odmiana algorytmu hMETIS, która umożliwia podział układu równocześnie na określoną liczbę części (ang. hMETIS-Kway). W algorytmie wprowadzony został nowy sposób łączenia modułów. Łączone są pary modułów, podobnie jak w zwykłym algorytmie hMETIS. Możliwe jest jednak kojarzenie par modułów, z których jeden został wcześniej skojarzony (ang. FirstChoice). Początkowy podział układu na najwyższym poziomie na określoną liczbę części, odbywa się z użyciem zwykłego algorytmu hMETIS. Podczas podziału na kolejnych, niższych poziomach zastosowano algorytm zachłanny. W tym celu, moduły są rozpatrywane kolejno w porządku losowym i mogą być przenoszone do sąsiednich części układu. Przeniesienie modułu jest możliwe w przypadku, gdy nie spowoduje to naruszenia ustalonego przydziału powierzchni dla poszczególnych części układu. Wybierana jest ta część układu, która zapewnia największą dodatnią zmianę kosztu połączeń. Zastosowanie algorytmu zachłannego zapewnia mniejsze nakłady obliczeniowe, w porównaniu do różnych modyfikacji algorytmu Fiduccia i Mattheysesa, stosowanych dla tego problemu [18].

Algorytm hMETIS umożliwia podział układów o bardzo dużych rozmiarach. Zapewnia otrzymanie rozwiązań o znacznie lepszej jakości, w porównaniu do innych algorytmów. Porównanie czasu potrzebnego na znalezienie rozwiązania również wypada na jego korzyść [17–18]. Algorytm hMETIS jest stosowany w programach rozmieszczania wykorzystujących algorytm min-cut: Dragon i Feng Shui. Jest również stosowany w programie rozmieszczania mPl, opartym na metodzie analitycznej oraz w programie MGP, który wykorzystuje algorytm symulowanego wyżarzania. Algorytm wielopoziomowego podziału hMETIS jest powszechnie dostępny w postaci plików wykonywalnych, dostępnych dla różnych systemów operacyjnych [19].

5. PROGRAMY ROZMIESZCZANIA WYKORZYSTUJĄCE ALGORYTM MIN-CUT

Programy rozmieszczania możemy podzielić na dwie grupy, w zależności od właściwości podłoża, na którym rozmieszczane są moduły. Jeżeli program rozmieszcza moduły na podłożu o ściśle określonych rozmiarach, z określonymi rozmiarami kanałów do prowadzenia połączeń oraz ustalonym sposobem prowadzenia zasilania, wówczas określamy go mianem programu typu fixed die (ang.). W przeciwnym przypadku jest programem typu variable-die (ang.), program ten stara się maksymalnie upakować moduły, w celu minimalizacji całkowitej powierzchni układu [20, 31–32]. Wśród obecnie stosowanych programów rozmieszczania, które wykorzystują algorytm min-cut, należy wymienić: Capo, Dragon, Feng Shui, QUAD [11, 20, 24–29, 31–34]. Programy te są przeznaczone przede wszystkim dla układów Standard Cell, w których moduły — komórki standardowe, rozmieszczane są w wierszach.

5.1. CAPO

Capo wykorzystuje hierarchiczny podział układu z wykorzystaniem algorytmu min-cut oraz algorytmu Fiduccia i Mattheysesa. Podział układu (części układu) zawierającego więcej niż 200 modułów wykonywany jest z użyciem wielopoziomowego algorytmu Fiduccia i Mattheysesa. „Płaski” (bez wielu poziomów) algorytm Fiduccia i Mattheysesa jest stosowany niezależnie do podziału układów zawierających 35–200 modułów oraz jest podstawą wielopoziomowego algorytmu podziału. Do podziału mniejszych części układu stosowany jest algorytm podziału i ograniczenia (ang. branch and bound). Układ jest dzielony liniami pionowymi oraz liniami poziomymi, które umożliwiają przypisanie modułów do konkretnego wiersza [20, 32–34].

W „płaskim” algorytmie Fiduccia i Mattheysesa zastosowano szereg modyfikacji. Moduły o dużej powierzchni często są przyczyną gorszych rezultatów, w przypadku tradycyjnej implementacji algorytmu. Moduły te są często na pierwszym miejscu wśród modułów przeznaczonych do przeniesienia, które z kolei może nie być wykonane ze względu na naruszenie ustalonego przydziału powierzchni dla powstających części. Mniejsze moduły, dające taką samą zmianę kosztu, nie są wtedy rozpatrywane, ze względu na oszczędność czasu. Rozwiązanie problemu polega na usunięciu zbyt dużych modułów, z pierwszych miejsc list modułów, dających taką samą zmianę kosztu. Ponadto, trwale ustalony przydział powierzchni dla powstających części, bez żadnej możliwości jego naruszenia, również jest przyczyną gorszych rezultatów. Rozwiązanie polega na uruchomieniu algorytmu kilka razy dla danego problemu podziału. Rezultaty podziału w danej próbie, są początkowym podziałem dla następnej. W pierwszej próbie tolerancja naruszenia przydziału powierzchni jest duża, w kolejnych krokach zmniejsza się do wartości pożądanej, która jest ściśle określana dla każdego podziału [20–22].

Modyfikacje wprowadzono również w wielopoziomowym algorytmie Fiduccia i Mattheysesa. Zastosowano inny sposób łączenia modułów. Łączone są pary modułów,

którym odpowiada największa sumaryczna waga. W tej sumie mają udział wszystkie węzły układu, do których należy dana para modułów. Udział węzłów posiadających dwie końcówki wynosi 2, dla pozostałych jest równy 1. W dodatku algorytm zapobiega łączeniu modułów, którym odpowiada bardzo duża sumaryczna waga. Łączenie modułów kończy się w momencie, gdy w wyniku jego zastosowania, układ składa się z nie więcej niż 200 modułów. Podział z użyciem wielopoziomowego algorytmu Fiduccia i Mattheysesa może być wykonywany niezależnie kilka razy dla danego problemu, a najlepsze rozwiązanie jest zapisywane. Wykonywany jest również dodatkowy podział z użyciem algorytmu wielopoziomowego podziału, w którym początkowe położenie modułów określa najlepsze rozwiązanie osiągnięte w poprzednich próbach. Taki podział określany jest terminem V-cycle (ang.). W programie Capo standardowo dwa razy są powtarzane następujące sekwencje: dwie niezależne próby podziału i jeden podział typu V-cycle [17, 20–21].

Algorytm podziału i ograniczenia jest stosowany do podziału części układu, zawierających mniej niż 35 modułów. Algorytm polega na przechodzeniu drzewa decyzyjnego, którego liście reprezentują wszystkie możliwe rozwiązania. Drzewo decyzyjne jest sprawdzane częściowo. W przypadku, gdy dolne ograniczenie częściowego rozwiązania nie jest mniejsze od kosztu najlepszego znalezionej rozwiązania, cała gałąź drzewa jest odcinana, ponieważ wszystkie jej liście reprezentują rozwiązanie o większym koszcie. Opisany algorytm ma określony czas na znalezienie rozwiązania. W przypadku przekroczenia limitu czasu, wykorzystywany jest „płaski” algorytm Fiduccia i Mattheysesa. Algorytm podziału i ograniczenia jest stosowany również do rozmieszczania modułów w częściach układu, które zawierają nie więcej niż 7-8 modułów. Moduły rozmieszczane są w jednym wierszu tak, aby nie zachodziły na siebie i długość połączeń była minimalna. Zastosowanie algorytmu podziału i ograniczenia dla podziału części układu oraz rozmieszczania modułów, zapewnia otrzymanie lepszych wyników. Czynności te składają się na etap rozmieszczania szczegółowego modułów [23].

Podczas podziału danej części układu, Capo uwzględnia połączenia modułów znajdujących się w danej części układu, z modułami znajdującymi się w pozostałych częściach układu. Każdemu „zewnętrznemu” modułowi, połączonemu z dowolnym modułem w dzielonej części układu, początkowo jest przypisane ustalone położenie w środku jego własnej części. Następnie, wszystkie „zewnętrzne” moduły znajdujące się po obydwu stronach linii podziału są łączone ze sobą, tworząc dwa „zewnętrzne” moduły po obydwu stronach linii podziału. Moduły te uwzględniają połączenia modułów dzielonej części układu, z modułami znajdującymi się w pozostałych częściach. Węzły, do których należą obydwa „zewnętrzne” moduły, po obydwu stronach linii podziału, nie są rozpatrywane, ponieważ przecięcie węzła linią podziału jest wtedy nieuniknione [23].

Capo jest programem typu fixed-die. Program równomiernie rozmieszcza wolne przestrzenie (ang. white space) w całym układzie, co przyczynia się do zmniejszenia skupienia połączeń. Jest algorytmem stochastycznym. Najlepsze rozwiązanie z pięciu

niezależnych prób rozmieszczenia jest zwykle znacznie lepsze od średniego kosztu rozwiązań. Capo pozwala na obecność makrobloków w układzie, które muszą mieć jednak ustaloną pozycję [20, 32, 34]. Program jest powszechnie dostępny w postaci kodu źródłowego oraz plików wykonywalnych [30].

5.2. DRAGON

Dragon wykorzystuje hierarchiczny podział układu z wykorzystaniem algorytmu min-cut oraz algorytmu hMETIS. Na każdym poziomie hierarchii podziału, poszczególne części układu są dzielone równocześnie na cztery ćwiartki z użyciem algorytmu hMETIS. Podział kończy się w momencie, gdy w każdej części układu znajduje się mniej niż 7 modułów [18, 24, 32–34].

Na końcu każdego poziomu hierarchii podziału, wszystkie części układu są zamieniane miejscami między sobą, w celu osiągnięcia minimalnej długości połączeń. Podczas podziału danej części układu na części, nie są uwzględniane połączenia z modułami znajdującymi się w pozostałych częściach układu. Powyższe rozwiązanie, w połączeniu z zamianą miejscami poszczególnych części układu na końcu każdego poziomu podziału, umożliwia uzyskanie lepszych wyników. Po zakończeniu podziału układu na części, następuje korekcja rozmieszczenia poszczególnych modułów, które jest określone przynależnością do poszczególnych części układu. Korekcja odbywa się z wykorzystaniem algorytmu symulowanego wyżarzania [24, 33–34].

Następnie jest wykonywany etap rozmieszczania szczegółowego modułów, podczas którego następuje dalsza korekcja ich położenia. Najpierw moduły są odpowiednio przesuwane, aby usunąć wszelkie zachodzenia modułów na siebie (ang. legalization), które były wcześniej dozwolone. Następnie odbywa się zamiana modułów parami, w celu osiągnięcia mniejszej długości połączeń. Metoda zamiany parami akceptuje jedynie zamiany położenia modułów, które prowadzą do zmniejszenia długości połączeń. Zamiana modułów odbywa się w ograniczonym sąsiedztwie, w kierunku pionowym lub poziomym. Po każdej zamianie modułów usuwane są zachodzenia modułów lub wolne przestrzenie w wierszach modułów, które mogły powstać wskutek zamiany. W tym celu może być korygowane położenie modułów, znajdujących się po prawej stronie zamienianych modułów, w obydwu wierszach [7, 24, 33, 35].

Dragon standardowo pakuje moduły od lewej do prawej strony układu, co jest właściwością programów typu variable-die. Posiada jednak również funkcję równomiernego rozmieszczania wolnych przestrzeni tak, jak w programach typu fixed-die. Funkcja ta umożliwia zmniejszenie skupienia połączeń, kosztem większej długości połączeń. Dragon jest algorytmem stochastycznym. Posiada funkcje, umożliwiające rozmieszczanie modułów z minimalizacją opóźnień występujących w układzie. Dragon nie umożliwia obecności makrobloków w układzie [32, 34, 36]. Program jest powszechnie dostępny w postaci plików wykonywalnych [30].

5.3. FENG SHUI

Feng Shui wykorzystuje hierarchiczny podział układu z wykorzystaniem algorytmu min-cut oraz algorytmu hMETIS. Układ jest dzielony na przemian liniami pionowymi oraz poziomymi, dopóki każda z części układu nie będzie zawierać dokładnie jednego modułu. Preferowany jest podział na dwie równe części. Linie poziome nie muszą pokrywać się z granicami wierszy [25–28, 31–32].

Podział układu odbywa się z wykorzystaniem algorytmu hMETIS. Podczas podziału danej części układu, początkowe położenie modułów dla tego algorytmu jest określone przez algorytm wstępnego podziału (ang. *iterative deletion*). Podczas podziału danej części układu są uwzględniane połączenia z modułami znajdującymi się w innych częściach układu. Stosowana jest iteracyjna korekcja podziału układu, podobnie jak w rozwiązaniu zaproponowanym przez Breuera, ale z uwzględnieniem połączeń z modułami znajdującymi się w pozostałych częściach układu — rozdział 3, rysunek 6. Podczas podziału układu liniami poziomymi w końcowym etapie podziału, często występuje sytuacja polegająca na tym, że daną część układu nie można podzielić na dwie części, zgodnie z granicą wierszy. Powodem tego jest brak możliwości podziału, bez naruszenia ustalonego przydziału powierzchni dla powstających części (ang. *narrow region problem*). Sytuacja ta może wystąpić nawet w przypadku, gdy powierzchnia dzielonej części układu jest większa od sumy powierzchni wszystkich modułów. Jedno z rozwiązań problemu polega na podziale danej części układu bez ustalonego przydziału powierzchni dla powstających części tak, aby koszt połączeń między nimi był minimalny. Moduły są następnie przesuwane i pakowane od lewej do prawej tak, aby nie występowało zachodzenie modułów na siebie oraz nie występowały wolne przestrzenie w układzie. Rozwiązanie to powoduje jednak znaczne różnice w długości wierszy oraz znaczny wzrost długości połączeń. W celu rozwiązania tego problemu, w programie Feng Shui zastosowano odpowiednie przesuwanie poziomych linii podziału (ang. *fractional cut*). Części układu są dzielone bez ustalonego przydziału powierzchni tak, aby koszt połączeń między nimi był minimalny. Następnie pozioma linia podziału jest przesuwana we właściwe miejsce, między powstałe części. Przesuwanie poziomych linii podziału, które mogą biegnąć w innych miejscach niż granica wierszy, powoduje, że moduły nie są przypisane do konkretnych wierszy. Ponadto, występuje zachodzenie modułów na siebie, ponieważ przesuwanie linii podziału nie gwarantuje, że moduły będą w całości położone wewnątrz własnych części. Powyższe rozwiązanie ostatecznie zapewnia jednak otrzymanie lepszych wyników, ze względu na równomierne rozmieszczenie modułów [25, 27–28].

Po wykonaniu podziału układu, konieczne jest ustalenie przydziału modułów do poszczególnych wierszy oraz usunięcie zachodzenia modułów na siebie. Czynność ta odbywa się kolejno dla wszystkich wierszy. Wybierana jest grupa modułów bez przydziału, posiadających największą współrzędną y w układzie tak, aby powierzchnia tych modułów była większa od powierzchni całego wiersza modułów. Następnie wybierane są moduły, które powinny być umieszczone w bieżącym wierszu oraz moduły

przeznaczone do następnego wiersza. Koszt przydziału modułu do bieżącego wiersza jest równy kwadratowi odległości między bieżącą i ostateczną pozycją modułu. Koszt przydziału do następnego wiersza jest równy kwadratowi odległości do następnego wiersza. Przydział modułów odbywa się z wykorzystaniem programowania dynamicznego. Następnie, moduły przydzielone do poszczególnych wierszy są pakowane od lewej do prawej bez zachodzenia na siebie, według współrzędnej x modułów. Powyższe czynności są powtarzane do momentu, gdy wszystkie moduły zostaną przydzielone do poszczególnych wierszy [25, 28].

Program Feng Shui stosuje również etap rozmieszczania szczegółowego modułów, w którym wykorzystywany jest algorytm podziału i ograniczenia. Celem korekcy rozmieszczenia jest osiągnięcia mniejszej długości połączeń dla małych grup modułów, liczących 6 lub mniej modułów. Korekcja rozmieszczenia odbywa się dla grup modułów położonych w jednym wierszu lub w sąsiednich wierszach [25].

Feng Shui jest programem typu variable-die i zawsze pakuje moduły od lewej do prawej strony układu. Program stara się zmieniać długość wierszy, w celu minimalizacji powierzchni układu. Feng Shui jest algorytmem stochastycznym. Feng Shui pozwala na obecności makrobloków w układzie [25, 31–32]. Program jest powszechnie dostępny w postaci plików wykonywalnych [30].

5.4. QUAD

QUAD jest programem wykorzystującym hierarchiczny podział układu z użyciem algorytmu min-cut. Na każdym poziomie hierarchii podziału, poszczególne części układu są dzielone równocześnie na cztery ćwiartki z użyciem wielopoziomowego algorytmu podziału, zaproponowanego przez Alperta i inn. [11, 16, 25, 29].

Koszt połączeń w programie QUAD może być równy sumie długości minimalnych drzew rozpinających, wyznaczanych dla wszystkich węzłów układu [2]. Minimalne drzewo rozpinające dla danego węzła, łączy wszystkie ćwiartki powstające podczas podziału, w których znajduje się przynajmniej jedna z końcówek węzła. Autorzy programu wychodzą z założenia, że funkcja kosztu oparta na minimalnym drzewie rozpinającym, umożliwia lepszą estymację kosztu połączeń, niezbędnych do połączenia wszystkich modułów w układzie. W programie możliwe jest również zastosowanie takiej samej funkcji kosztu, jak w tradycyjnym algorytmie min-cut, nie wymaga to dużych zmian. QUAD pozwala na stosowanie różnych wag połączeń dla kierunku poziomego i pionowego, co umożliwia wprowadzenie preferencji dla wybranego kierunku połączeń. Uwzględniane są połączenia z modułami znajdującymi się poza dzieloną częścią układu, w wyniku zastosowania modyfikacji Dunlopa i Kernighana. Podczas podziału danej części, stosowana jest iteracyjna korekcja jej podziału. Korekta jest wykonywana cyklicznie dla wszystkich ćwiartek dzielonej części, aż do momentu ustalenia wyniku. Stosowana jest również iteracyjna korekcja podziału, umożliwiająca zmianę przydziału modułów do danej ćwiartki, który został wcześniej określony podczas poprzedniego

poziomu hierarchii podziału. QUAD jest zaliczany do programów rozmieszczania globalnego [29].

6. PODSUMOWANIE

Niniejsza praca jest drugą częścią przeglądu, poświęconego metodom rozmieszczania w projektowaniu topografii układów VLSI. Zawiera szczegółowy opis algorytmu min-cut oraz programów rozmieszczania, wykorzystujących ten algorytm.

Chociaż algorytm min-cut powstał blisko 30 lat temu, jego modyfikacje nadal należą do najpowszechniej stosowanych metod. Umożliwiają rozmieszczenie modułów nawet w przypadku układów o bardzo dużych rozmiarach. Ostatnie lata wskazują na wzrastające zainteresowanie badaczy tym algorytmem. Powstało wiele programów rozmieszczania dla układów Standard Cell, których podstawą jest ten właśnie algorytm. We współczesnych programach dla układów Standard Cell (Capo, Dragon, Feng Shui), pojawiły się nowe tendencje. Algorytm min-cut wykorzystywany jest do rozmieszczania globalnego. Rozmieszczanie szczegółowe wykonywane jest z użyciem innych metod. Dalsze badania skupione są na efektywnym rozmieszczaniu makrobloków w układach o topografii mieszanej.

Algorytm min-cut jest nadal bardzo efektywnym narzędziem rozmieszczania modułów. Należy pamiętać o konieczności modyfikacji wraz z rosnącymi potrzebami projektowania coraz to większych systemów w jednym module scalonym.

7. BIBLIOGRAFIA

1. M.J.S. Smith: *Application-Specific Integrated Circuits*, Addison Wesley Longman, 1997.
2. A. Kos, Z. Nagórny: *Estymacja długości połączeń w układach VLSI*, IV Krajowa Konferencja Elektroniki, Darłówko Wschodnie, czerwiec 2005, ss. 159-164.
3. B.T. Preas, M. J. Lorenzetti (red.): *Physical Design Automation of VLSI Systems*, Menlo Park, Benjamin-Cummings, 1988
4. T. Ohtsuki (red.): *Layout Design and Verification*, Elsevier Science Publishers B. V. (North Holland), 1986.
5. M.M. Vai: *VLSI Design*, CRC Press, 2001.
6. C. Sechen: *VLSI Placement and Global Routing Using Simulated Annealing*, Boston, Kluwer Academic Publishers, 1988.
7. M.A. Breuer (red.): *Automatyczne projektowanie maszyn cyfrowych*, Warszawa, PWN, 1976.
8. W. Wolf: *Modern VLSI Design: a systems approach*, Englewood Cliffs, New Jersey, PTR Prentice Hall, 1994.
9. K. Shahookar, P. Mazumder: *VLSI Cell Placement Techniques*, ACM Computing Surveys, 1991, vol. 23, pp. 143-220.
10. T.C. Hu, E.S. Kuh (red.): *VLSI Circuit Layout: Theory and Design*, New York, IEEE Press, 1985.
11. C. J. Alpert, G.J. Nam, P. G. Villarrubia: *Effective Free Space Management for Cut-Based Placement via Analytical Constraint Generation*, IEEE Transactions on Computer Aided Design, 2003, vol. 22, pp. 1343-1353.

12. B. Krishnamurthy: *An Improved Min-Cut Algorithm for Partitioning VLSI Networks*, IEEE Transactions on Computers, 1984, vol. 33, pp. 438–446.
13. A.E. Dunlop, B.W. Kernighan: *A Procedure for Placement of Standard-Cell VLSI Circuits*, IEEE Transactions on Computer Aided Design, 1985, vol. 4, pp. 92–98.
14. P.R. Suaris, G. Kedem: *An Algorithm for Quadrissection and Its Application to Standard Cell Placement*, IEEE Transactions on Circuits and Systems, 1988, vol. 35, pp. 294–302.
15. A.E. Caldwell, A.B. Kahng, I.L. Markov: *Design and Implementation of Move-Based Heuristics for VLSI Hypergraph Partitioning*, ACM Journal of Experimental Algorithmics, 2000, vol. 5, artykuł 5, <http://www.eecs.umich.edu/~imarkov/pubs/jour/j004.pdf>.
16. C.J. Alpert, J.H. Huang, A.B. Kahng: *Multilevel Circuit Partitioning*, IEEE Transactions on Computer Aided Design, 1998, vol. 17, pp. 655–667.
17. G. Karypis, R. Aggarwal, V. Kumar, S. Shekhar: *Multilevel Hypergraph Partitioning: Applications in VLSI Domain* IEEE-Transactions on VLSI Systems, 1999, vol. 7, pp. 69–79.
18. G. Karypis, V. Kumar: *Multilevel k-way Hypergraph Partitioning*, Design Automation Conference, 1999, pp. 343–348, <http://www.users.cs.umn.edu/~karypis/publications/Papers/PDF/khmetis.pdf>.
19. University of Minnesota, USA; George Karypis's Home Page, <http://www.users.cs.umn.edu/~karypis/metis/hmetis/download.html>.
20. A.E. Caldwell, A.B. Kahng, I.L. Markov: *Can Recursive Bisection Alone Produce Routable Placements?*, Design Automation Conference, 2000, pp. 477–482, <http://www.eecs.umich.edu/~imarkov/pubs/conf/c015.pdf>.
21. A.E. Caldwell, A.B. Kahng, I.L. Markov: *Improved Algorithms for Hypergraph Bipartitioning*, Asia and South Pacific Design Automation Conference, 2000, pp. 661–666, <http://www.eecs.umich.edu/~imarkov/pubs/conf/c013.pdf>.
22. A.E. Caldwell, A.B. Kahng, I.L. Markov: *Iterative Partitioning with Varying Node Weights*, VLSI Design, 2000, vol. 11, pp. 249–258, <http://www.eecs.umich.edu/~imarkov/pubs/jour/j006.pdf>.
23. A.E. Caldwell, A.B. Kahng, I.L. Markov: *Optimal Partitioners and End-Case Placers for Standard-Cell Layout*, IEEE Transactions on Computer Aided Design, 2000, vol. 19, pp. 1304–1313.
24. M. Wang, X. Yang, M. Sarrafzadeh: *Dragon2000: Standard-Cell Placement Tool for Large Industry Circuits*, International Conference on Computer Aided Design, 2000, pp. 260–263, <http://er.cs.ucla.edu/Dragon/papers/dragon.pdf>.
25. A.R. Agnihotri, S. Ono, C. Li, M.C. Yildiz, A. Khatkhate, C.K. Koh, P.H. Madden: *Mixed Block Placement via Fractional Cut Recursive Bisection*, IEEE Transactions on Computer Aided Design, 2005, vol. 24, pp. 748–761.
26. M.C. Yildiz, P.H. Madden: *Improved Cut Sequences for Partitioning Based Placement*, Design Automation Conference, 2001, pp. 776–779, <http://vlsicad.cs.binghamton.edu/pubs/Yildiz01dac.pdf>.
27. M.C. Yildiz, P.H. Madden: *Global Objectives for Standard-Cell Placement*, Great Lakes Symposium on VLSI, 2001, pp. 68–72, <http://vlsicad.cs.binghamton.edu/pubs/Yildiz010068.pdf>.
28. A. Agnihotri, M.C. Yildiz, A. Khatkhate, A. Mathur, S. Ono, P.H. Madden: *Fractional Cut: Improved Recursive Bisection Placement*, International Conference on Computer Aided Design, 2003, pp. 307–310, <http://vlsicad.cs.binghamton.edu/pubs/Agnihotri03.pdf>.
29. D.J.H. Huang, A.B. Kahng: *Partitioning-Based Standard-Cell Global Placement with an Exact Objective*, International Symposium on Physical Design, 1997, pp. 18–25, <http://citeseer.ist.psu.edu/71921.html>.
30. VLSI CAD Bookshelf Slots and Entries, A.B. Kahng, I.L. Markov, <http://vlsicad.eecs.umich.edu/BK/Slots/slots/WirelengthdrivenStandardCellPlacement.html>.
31. P.H. Madden: *Reporting of Standard Cell Placement Results*, IEEE Transactions on Computer Aided Design, 2002, vol. 21, pp. 240–247.
32. S.N. Adya, M.C. Yildiz, I.L. Markov, P.G. Villarrubia, P.N. Parakh, P.H. Madden: *Benchmarking for Large-Scale Placement and Beyond*, IEEE Transactions on Computer Aided Design, 2004, vol. 23, pp. 472–486.

33. O. Liu, M. Marek-Sadowska: *A Study of Netlist Structure and Placement Efficiency*, IEEE Transactions on Computer Aided Design, 2005, vol. 24, pp. 762-772.
34. C.C. Chang, J. Cong, M. Romesis, M. Xie: *Optimality and Scalability Study of Existing Placement Algorithms*, IEEE Transactions on Computer Aided Design, 2004, vol. 23, pp. 537-548.
35. M. Gajęcki, A. Kos: *A Problem of Optimization of Topography of VLSI Circuit*, Proc. of the XIXth National Conference on Circuit Theory and Electronic Networks, Kraków-Krynica (Poland), October 1996, pp. II/307-312.
36. X. Yang, B.K. Choi, M. Sarrafzadeh: *Timing-Driven Placement using Design Hierarchy Guided Constraint Generation*, International Conference on Computer Aided Design, 2002, pp. 177-184
<http://er.cs.ucla.edu/Dragon/papers/timingdragon.pdf>.

Z. NAGÓRNY, A. KOS

THE DESIGN OF THE VLSI CIRCUIT LAYOUT – PART 2. MIN-CUT ALGORITHM

Summary

The design process of the VLSI circuits requires the use of computer aided design tools. This paper is the second part of the survey of the cell placement techniques for digital VLSI circuits. In this part of the survey, the min-cut algorithm is presented. The Kernighan-Lin algorithm and its modifications, which are the base of the min-cut algorithm are described. Then, the Fiduccia-Mattheyses algorithm is described. The computation time of the Fiduccia-Mattheyses algorithm increases only slightly more than linearly with the number of logic cells in the circuit. It is a very important improvement. Some modifications of the min-cut algorithm are presented. The terminal propagation and the quadrisection algorithm are described. The application of the min-cut algorithm for the building block design style is presented. The principles of the multilevel circuit partitioning algorithm are described. Two multilevel circuit partitioning algorithms are characterized: hMETIS and hMETIS-Kway. Nowadays the tools used for the cell placement, which utilize the presented algorithms are characterized: Capo, Dragon, Feng Shui, QUAD. Some conclusions concerning described techniques and tools are presented.

Keywords: VLSI, physical design process, integrated circuit, layout, computer aided design, Standard Cell, cell placement, circuit partitioning, Kernighan-Lin algorithm, Fiduccia-Mattheyses algorithm, min-cut algorithm, multilevel circuit partitioning, hMETIS algorithm, Capo, Dragon, Feng Shui, QUAD

INFORMACJE DLA AUTORÓW

Redakcja przyjmuje do publikowania prace oryginalne, przeglądowe i monograficzne wchodzące w zakres szeroko pojętej elektroniki. Ponieważ KWARTALNIK ELEKTRONIKI I TELEKOMUNIKACJI jest czasopismem Komitetu Elektroniki i Telekomunikacji Polskiej Akademii Nauk, w związku z tym na jego łamach znajdują się prace naukowe dotyczące podstaw teoretycznych i zastosowań z zakresu elektroniki, telekomunikacji, mikroelektroniki, optoelektroniki, radiotechniki i elektroniki medycznej.

Artykuły powinny charakteryzować oryginalne ujęcie zagadnienia, własna klasyfikacja, krytyczna ocena (teorii lub metod), omówienie aktualnego stanu, lub postępu danej gałęzi techniki oraz omówienie perspektyw rozwojowych. Artykuły publikowane w innych czasopismach nie mogą być kierowane do druku w Kwartalniku Elektroniki i Telekomunikacji w drugiej kolejności zgłoszenia.

Objętość artykułu nie powinna przekraczać 30 stron po około 1800 znaków na stronie, w tym rysunki i tabele.

Wymagania podstawowe.

Artykuły należy nadsyłać na wyraźnym, jednostronnym, czarno-białym wydruku komputerowym. Wydruk w formacie A4 powinien mieć znormalizowaną liczbę wierszy i znaków w wierszu (30 wierszy po 60 znaków w wierszu), w dwóch egzemplarzach, w języku polskim lub angielskim wybranym przez autora. Do wydruku powinna być dołączona dyskietka z elektronicznym tekstem artykułu. Preferowane edytory to WORD 6 lub 8. Układ artykułu (w wersji podstawowej) musi być następujący:

- Tytuł.
- Autor (imię i nazwisko autora/ów).
- Miejsce pracy (nazwa instytucji, miejscowość, adres. + ew. adres elektroniczny (e-mail)).
- Zwięzłe streszczenie powinno być w języku takim, w jakim jest pisany artykuł (wraz ze słowami kluczowymi).
- Tekst podstawowy powinien mieć następujący układ:

1. WPROWADZENIE
2. np. TEORIA
3. np. WYNIKI NUMERYCZNE
- 3.1.
- 3.2.
4.
5.
6. PODSUMOWANIE
7. ew. PODZIĘKOWANIA
8. BIBLIOGRAFIA

- Układ streszczenia w dodatkowej wersji językowej powinien być następujący:

AUTOR (inicjał imienia i nazwisko).

TYTUŁ (w języku angielskim – o ile artykuł pisany jest w języku polskim i na odrwót).

Obszerne do 3600 znaków streszczenia (wraz z słowami kluczowymi) w języku:

- a. angielskim, gdy artykuł pisany jest w języku polskim.
- b. polskim, gdy artykuł pisany jest w języku angielskim.

Streszczenie to powinno pozwolić czytelnikowi na uzyskanie istotnych informacji zawartych w pracy. Z tego względu w streszczeniu tym mogą być cytowane numery istotnych wzorów, rysunków i tabel zawartych w podstawowej wersji językowej.

- Wszystkie strony muszą mieć numerację ciągłą.

Sposób pisania tekstu.

Tekst powinien być pisany bez używania wyróżnień, a w szczególności nie dopuszcza się spacjiowania, podkreślania i pisania tekstu dużymi literami z wyjątkiem wyrazów, które umownie pisze się dużymi literami (np. FORTRAN). Proponowane wyróżnienia Autor może zaznaczyć w maszynopisie zwykłym ołówkiem za pomocą przyjętych znaków adjustacyjnych, np. podkreślenie linią przerywaną oznacza spacjiowanie (rozstrzeżenie), podkreślenie linią ciągłą – pogrubienie, podkreślenie wężykiem — kursywa.

Tekst powinien być napisany z podwójnym odstępem między wierszami, tytuły i podtytuły małymi literami. Marginesy z każdej strony powinny mieć około 35 mm. Wielkość czcionki wydruku powinna być zbliżona co najmniej co wielkości czcionki maszyny do pisania (minimum 12 punktów). Przy podziale pracy na rozdziały i podrozdziały cyfrowe ich oznaczenia nie powinny być większe niż II stopnia (np. 4.1.1.).

Sposób pisania tabel.

Tabele powinny być pisane na oddzielnych stronach. Tytuły rubryk pionowych i poziomych powinny być napisane małymi literami z podwójnym odstępem między wierszami. Przypisy (notki) dotyczące tabel należy pisać bezpośrednio pod tabelami. Tabele należy numerować kolejno liczbami arabskimi, u góry każdej tabeli podać tytuł dwujęzyczny. W pierwszej kolejności w podstawowej wersji językowej, a później w dodatkowej wersji językowej. Tabele umieścić na końcu maszynopisu. Przyjmowane są tabele algorytmów i programy na wydrukach komputerowych. W tym przypadku zachowany jest ich oryginalny układ. Tabele powinny być cytowane w tekście.

Sposób pisania wzorów matematycznych.

Rozmieszczenie znaków, cyfr, liter i odstępów powinno być zbliżone do rozmieszczenia elementów druku. Wskaźniki i wykładniki potęg powinny być napisane wyraźnie i być prawidłowo obniżone lub podwyższone w stosunku do linii wiersza podstawowego. Znaki nad literami i cyframi, całkami i in. symbolami (strzałki, linie, kropki, daszki) powinny być umieszczone dokładnie nad tymi elementami, do których się odnoszą. Numery wzorów cyframi arabskimi powinny być kolejne i umieszczone w nawiasach okrągłych z prawej strony. Nazwy jednostek, symbole literowe i graficzne powinny być zgodne z wytycznymi IEE (International Electronical Commision) oraz ISO (International Organization of Standarization).

Powołania.

Powołania na publikacje powinny być umieszczone na ostatnich stronach tekstu pod tytułem „Bibliografia”, opatrzone numeracją kolejną bez nawiasów. Numeracja ta powinna być zgodna z odnośnikami w tekście artykułu. Przykłady opisu publikacji:

- periodycznej 1. F. Valdoni: A new milimetre wave satelite. E.T.T. 1990, vol. 2, no 5, pp. 141–148
- nieperiodyczne 2. K. Andersen: A resource allocation framework. XVI International Symposium, Stockholm (Sweden), may 1991, paper A 2,4
- książki 3. Y.P. Tvidis: Operation and modeling of the MOS transistors. New York, McGraw-Hill, 1987, p. 553

Materiały ilustracyjne.

Rysunki powinny być wykonane wyraźnie, na papierze gładkim lub milimetrowym w formacie nie mniejszym niż 9 × 12 cm. Mogą być także w postaci wydruku komputerowego (preferowany edytor Corel Draw). Fotografie lub diapozytywy przyjmowane są raczej czarno-białe w formacie nie przekraczającym 10 × 15 cm. Na marginesie każdego rysunku i na odwrocie fotografii powinno być napisane ołówkiem imię i nazwisko Autora oraz skrót tytułu artykułu, do którego są przeznaczone oraz numer rysunku. Spis podpisów pod rysunki i fotografie powinny być umieszczone na oddzielnej stronie. Podpisy pod rysunkami (fotografiami) powinny być dwujęzyczne: w pierwszej kolejności w podstawowej wersji językowej, a później w dodatkowej wersji językowej. Rysunki powinny być cytowane w tekście.

Uwagi końcowe.

Na odrębnej stronie powinny być podane następujące informacje:

- adres do korespondencji z kodem pocztowym (domowy lub do miejsca pracy),
- telefon domowy i/lub do miejsca pracy,
- adres e-mailowy (jeśli autor posiada).

Autorowi przysługuje bezpłatnie 20 odbitek artykułu. Dodatkowe egzemplarze odbitek, lub cały zeszyt Autor może zamówić u wydawcy na własny koszt.

Autora obowiązują korekta autorska, którą powinien wykonać w ciągu 3 dni od daty otrzymania tekstu z Redakcji oraz zwrócić osobiście lub listownie pod adres Redakcji. Korekta powinna być naniesiona na przekazanych Autorowi szpaltach na marginesach ew. na osobnym arkuszu w przypadku uzupełnień tekstu większych niż dwa wiersze. W przypadku nie zwrócenia korekty w terminie, korektę przeprowadza Redakcja Techniczna Wydawcy. Redakcja prosi Autorów o powiadomienie ją o zmianie miejsca pracy i adresu prywatnego.

INFORMATION FOR AUTHORS OF K.E.T.

The editorial staff will accept for publishing only original monographic and survey papers concerning widely understood electronics. Because of the fact that KWARTALNIK ELEKTRONIKI I TELEKOMUNIKACJI is a journal of the Committee for Electronics and Telecommunications of Polish Academy of Science, it presents scientific works concerning theoretical bases and applications from the field of electronics, telecommunications, microelectronics, optoelectronics, radioelectronics and medical electronics.

Articles should be characterised by original depiction of a problem, its own classification, critical opinion (concerning theories or methods), discussion of an actual state or a progress of a given branch of a technique and discussion of development perspectives.

An article published in other magazines can not be submitted for publishing in K.E.T. The size of an article can not exceed 30 pages, 1800 character each, including figures and tables.

Basic requirements

The article should be submitted to the editorial staff as a one side, clear, black and white computer printout in two copies. The article should be prepared in English or Polish. Floppy disc with an electronic version of the article should be enclosed. Preferred wordprocessors: WORD 6 or 8.

Layout of the article.

- Title.
- Author (first name and surname of author/authors).
- Workplace (institution, adress and e-mail).
- Concise summary in a language article is prepared in (with keywords).
- Main text with following layout:
 - Introduction
 - Theory (if applicable)
 - Numerical results (if applicable)
 - Paragraph 1
 - Paragraph 2
 -
 -
 - Conclusions
 - Acknowledgements (if applicable)
 - References
- Summary in additional language:
 - Author (firs name initials and surname)
 - Title (in Polish, if article was prepared in English and vice versa)
 - Extensive summary, hawever not exceeching 3600 characters (along with keywords) in Polish, if article was prepared in English and vice versa). The summary should be prepared in a way allowing a reader to obtaoin essential information contained in the artide. For that reason in the summary author can place numbers of essential formulas, figures and tables from the article.

Pages should have continues numbering.

Main text

Main text can not contain formatting such as spacing, underlining, words written in capital letters (except words that are commonly written in capital letters). Author can mark suggested formatting with pencil on the margin of the article using commonly accepted adjusting marks.

Text should be written with double line spacing with 35 mro left and right margin. Titles and subtitles should be written with small letters. Titles and subtitles should be numberd using no mor than 3 levels (i.e. 4.1.1.).

Tables

Tables with their titles should be places on separate page at the end of the article. Titles of rows and columns should be written in small letters with double line spacing. Annotations concerning tables

should be placed directly below the table. Tables should be numbered with Arabic numbers on the top of each table. Table can consist algorithm and program listings. In such cases original layout of the table will be preserved. Table should be cited in the text.

Mathematical formulas

Characters, numbers, letters and spacing of the formula should be adequate to layout of main text. Indexes should be properly lowered or raised above the basic line and clearly written. Special characters such as lines, arrows, dots should be placed exactly over symbols which they are attributed to. Formulas should be numbered with Arabic numbers placed in brackets on the right side of the page. Units of measure, letter and graphic symbols should be printed according to requirements of IEE (International Electrotechnical Commission) and ISO (International Organisation of Standardisation).

References

References should be placed at the end of the main text with the subtitle „References“. References should be numbered (without brackets) adequately to references placed in the text. Examples of periodical [2], non-periodical [2] and book [3] references:

1. F. Valdoni: A new millimetre wave satellite. E.T.T. 1990, vol. 2, no 5, pp. 141–148
2. K. Anderson: A resource allocation framework. XVI International Symposium (Sweden). May 1991. paper A 2.4
3. Y.P. Tividis: Operation and modeling of the MOS transistors. New York. McGraw-Hill. 1987. p. 553

Figures

Figures should be clearly drawn on plain or millimetre paper in the format not smaller than 9×12 cm. Figures can be also printed (preferred editor – CorelDRAW). Photos or diapositives will be accepted in black and white format not greater than 10×15 cm. On the margin of each drawing and on the back side of each photo author's name and abbreviation of the title of article should be placed. Figure's captions should be given in two languages (first in the language the article is written in and then in additional language). Figure's captions should be also listed on separate page. Figures should be cited in the text.

Additional information

On the separate page following information should be placed:

- mailing address (home or office),
- phone (home or/and office),
- e-mail.

Author is entitled to free of charge 20 copies of article. Additional copies or the whole magazine can be ordered at publisher at the author's expense.

Author is obliged to perform the author's correction, which should be accomplished within 3 days starting from the date of receiving of the text from the editorial staff. Corrected text should be returned to the editorial staff personally or by mail. Correction marks should be placed on the margin of copies received from the editorial staff or if needed on separate pages. In the case when the correction is not returned in said time limit, correction will be performed by technical editorial staff of the publisher.

In case of changing of workplace or home address Authors are asked to inform the editorial staff.